

Ralf Gerkmann
Mathematisches Institut der
Ludwig-Maximilians-Universität München

Wintersemester 2024-25 / Sommersemester 2025

Lineare Algebra

Inhaltsverzeichnis

§ 1. Aussagenlogik	3
§ 2. Mengenlehre und Prädikatenlogik	10
§ 3. Relationen	19
§ 4. Abbildungen und Mächtigkeiten	30
§ 5. Algebraische Grundstrukturen und Matrizen	47
§ 6. Vektorräume, lineare Abbildungen und lineare Gleichungssysteme	62
§ 7. Die Lösung linearer Gleichungssysteme	74
§ 8. Linearkombinationen und lineare Unabhängigkeit	90
§ 9. Basen eines Vektorraums, Dimensionsbegriff	99
Literaturverzeichnis	108

§ 1. Aussagenlogik

Inhaltsübersicht

Unter einer **Aussage** verstehen wir einen (sprachlich oder in mathematischer Notation formulierten) Satz, von dem auf sinnvolle und objektive Weise gesagt werden kann, dass er **wahr** oder **falsch** ist. Mit Hilfe von logischen Symbolen $\neg, \wedge, \Rightarrow$ usw. lassen sich einfache Aussagen zu komplexeren Aussagen zusammensetzen. Die **Tautologien** bilden eine besonders wichtige Klasse zusammengesetzter Aussagen, weil sie für logische Schlüsse verwendet werden können. Aus solchen Schlüssen wiederum werden mathematische **Beweise** aufgebaut.

Wichtige Begriffe und Sätze

- Aussagen und ihre Wahrheitswerte (wahr oder falsch)
- Aussagenschema, Parameter
- Verknüpfung von Aussagen (Konjunktion, Disjunktion, Negation, Implikation, Äquivalenz)
- Tautologien und logische Schlüsse

(1.1) Definition Unter einer **Aussage** verstehen wir einen (sprachlich oder in mathematischer Notation formulierten) Satz, von dem auf sinnvolle und objektive Weise gesagt werden kann, dass er **wahr** oder **falsch** ist.

Die folgenden Sätze sind zweifellos Aussagen.

- (i) Heute ist Dienstag. (*wahr*, jedenfalls am 15.10.2024)
- (ii) $1 + 1 = 2$ (*wahr*)
- (iii) Es gibt eine natürliche Zahl, die größer ist als alle anderen natürlichen Zahlen.
(*Falsch*. Nehmen wir an, n wäre eine solche Zahl. Dann müsste $n > n + 1$ gelten.
Wir wissen aber, dass $n < n + 1$ gilt.)
- (iv) Die Summe der Innenwinkel eines beliebigen Dreiecks beträgt 180° .
(*wahr*, zumindest in der „normalen“ euklidischen Geometrie)
- (v) Jede differenzierbare Funktion ist stetig. (*wahr*)
- (vi) Jede gerade Zahl größer als zwei kann als Summe von zwei Primzahlen dargestellt werden.
(Dies ist die sog. *Goldbachsche Vermutung*. Zur Zeit ist noch unbekannt, ob sie wahr oder falsch ist.)

Dagegen sind die folgenden Sätze mit Sicherheit nicht als Aussagen zu bezeichnen.

- (i) Hallo!
- (ii) Mach endlich Deine Hausaufgaben!
- (iii) 10^{100} ist eine große Zahl
- (iv) Die Kreiszahl π ist ungefähr gleich 3.14.
- (v) $x^3 - 3x^2 - 3x + 1$
- (vi) $a^2 + b^2 = c^2$

Offenbar ist es sinnlos, einer Begrüßung oder einer Aufforderung einen Wahrheitswert zuzuordnen. Die Sätze (iii) und (iv) sind für eine Aussage nicht hinreichend objektiv. Der Ausdruck (v) ist ein **Term** (genauer gesagt, ein Polynom), für den die Feststellung *wahr* oder *falsch* ebenfalls keinen Sinn macht.

Satz (vi) ist für sich genommen keine Aussage, solange den Symbolen a , b und c keine Bedeutung zugeordnet wird. Legt man fest, dass $a = 3$, $b = 4$ und $c = 5$ sein soll, erhält man eine wahre Aussage, denn es gilt $3^2 + 4^2 = 9 + 16 = 25 = 5^2$. Für viele andere Belegungen von a, b, c (zum Beispiel $a = 1$, $b = 2$, $c = 3$) erhält man dagegen eine falsche Aussage. Legt man fest, dass a, b, c Seitenlängen eines rechtwinkligen Dreiecks sein sollen, wobei c der längsten Seite zugeordnet ist, dann erhält man wiederum eine wahre Aussage. Einer der häufigsten Fehler bei der Formulierung mathematischer Aussagen besteht darin, dass Bezeichnungen (wie hier a, b, c) verwendet werden, die zuvor nicht definiert wurden!

Die Gleichung $a^2 + b^2 = c^2$ ist also keine Aussage; statt dessen fällt sie eine allgemeinere Kategorie von Sätzen, die man unter dem Begriff „**Aussagenschema**“ zusammenfasst. Bei einem Aussagenschema handelt es sich um einen Satz, in dem eine Reihe von **Parametern** $x, y, \dots$ vorkommen, und der zu einer Aussage wird, wenn man die Parameter durch geeignete mathematische Objekte ersetzt. Beispielsweise wird $a^2 + b^2 = c^2$ zu einer Aussage, wenn man für a, b, c die Längen der Katheten und der Hypotenuse eines rechtwinkligen Dreiecks einsetzt. Auch der (sprachlich formulierte) Satz

„Die Zahl x ist eine Primzahl.“

ist ein Aussagenschema mit x als Parameter. Setzt man für x die Werte 4 oder 6 ein, so erhält man eine falsche Aussage. Setzt man dagegen 2 oder 13 ein, dann erhält man eine wahre Aussage.

Zu beachten ist, dass im Allgemeinen natürlich nicht jede Einsetzung eine **sinnvolle** Aussage liefert. Zum Beispiel würde es keinen Sinn machen, in der Gleichung $a^2 + b^2 = c^2$ für a die leere Menge $\emptyset$ einzusetzen, da nicht ohne Weiteres klar ist, was der Ausdruck $\emptyset^2 + b^2 = c^2$ bedeuten soll.

Einfache Aussagen können umgangssprachlich, zum Beispiel durch Bindewörter wie „und“, „oder“, oder auch durch bestimmte Symbole ($\vee$, $\wedge$) zu komplexeren Aussagen **verknüpft** werden. Der Wahrheitswert der neuen Aussage ist dann durch die Wahrheitswerte der verknüpften Aussagen festgelegt. Wie diese Festlegung im einzelnen aussieht, kann am einfachsten durch sog. **Wahrheitstabellen** beschrieben werden. Seien φ und ψ zwei Aussagen. Die folgenden Verknüpfungen von Aussagen sind in der Mathematik allgemein gebräuchlich.

- (i) **Konjunktion** $\varphi \wedge \psi$ „Es gilt φ und ψ .“

φ	ψ	$\varphi \wedge \psi$
w	w	w
w	f	f
f	w	f
f	f	f

Die erste Zeile der Tabelle bedeutet ausformuliert: „Sind die Aussagen φ und ψ beide wahr, dann ist auch die zusammengesetzte Aussage $\varphi \wedge \psi$ eine wahre Aussage.“ Beispielsweise ist der Satz

„Heute ist Mittwoch, und es gilt $1 + 1 = 2$.“

eine wahre Aussage - über den Erkenntniswert kann man geteilter Meinung sein. Wichtig hierbei ist, dass auch die zusammengesetzten Aussage entweder *wahr* oder *falsch* ist; in der mathematischen Logik ist kein Platz für „Halbwahrheiten“. So ist der Satz

„Heute ist Mittwoch, und es gilt $1 + 1 = 3$.“

auch am Mittwoch, dem 16. Oktober 2024 auf Grund des Eintrags in der zweiten Tabellenzeile eindeutig als *falsch* zu bezeichnen. (Am 18. Oktober 2024 entnimmt man der vierten Tabellenzeile, dass die Aussage *falsch* ist, denn in diesem Fall sind beide Teilaussagen falsch.)

- (ii) **Disjunktion** $\varphi \vee \psi$ „Es gilt φ oder ψ .“

φ	ψ	$\varphi \vee \psi$
w	w	w
w	f	w
f	w	w
f	f	f

Zum Beispiel ist die Aussage „Es gilt $1 + 2 = 3$ oder $3 + 5 = 7$.“ wahr (zweite Tabellenzeile). Ebenso stimmt für jede reelle Zahl a die Aussage „Es gilt $a \geq 0$ oder $a \leq 0$.“, und zwar unabhängig davon, welche konkrete Zahl a man dort einsetzt. Hier kommt zum Beispiel für $a = 0$ die erste, für $a = -2$ die dritte Zeile zur Anwendung.

Zu beachten ist, dass sich beim mathematischen „oder“ die beiden Aussagen φ und ψ nicht gegenseitig ausschließen, wie dies beim umgangssprachlichen „entweder - oder“ der Fall ist: Die Aussage $\varphi \vee \psi$ ist auch dann wahr, wenn die Aussagen φ und ψ beide zutreffen!

- (iii) **Negation** $\neg\varphi$ „ φ gilt nicht.“ / „ φ ist falsch.“

φ	$\neg\varphi$
w	f
f	w

Beispielsweise ist der Satz „Die Gleichung $1 + 1 = 3$ gilt nicht.“ eine wahre Aussage (laut zweiter Tabellenzeile), und der Satz „Die Gleichung $1 + 1 = 2$ gilt nicht.“ ist falsch (laut erster Zeile).

- (iv) **Implikation** $\varphi \Rightarrow \psi$ „Aus φ folgt ψ .“ / „Wenn φ gilt, dann gilt auch ψ .“ / „ φ ist eine **hinreichende** Bedingung für ψ .“ / „ ψ ist eine **notwendige** Bedingung für φ .“

φ	ψ	$\varphi \Rightarrow \psi$
w	w	w
w	f	f
f	w	w
f	f	w

Man bezeichnet φ als die **Prämisse**, ψ als die **Konklusion** der Implikation $\varphi \Rightarrow \psi$. Bemerkenswert ist die Festlegung in der vierten Zeile: Wenn die Prämisse falsch ist, dann gilt die Implikation $\varphi \Rightarrow \psi$ auf jeden Fall als wahr, unabhängig vom Wahrheitswert der Aussage Konklusion. So gesehen ist

„Wenn $1 + 1 = 3$ ist, dann gilt auch $2 + 7 = 11$.“

eine wahre (wenn auch nicht besonders nützliche) Aussage. Logiker verwenden dafür den Ausspruch „*Ex falso quodlibet*“, d.h. aus etwas Falschem folgt alles Mögliche.

Bei der Implikation ist zu beachten, dass es zwischen den Aussagen A und B kein *kausaler* Zusammenhang bestehen muss, damit die Implikation $A \Rightarrow B$ zu einer wahren Aussage wird. Es kommt nur auf die Wahrheitswerte von φ und ψ an. Beispielsweise ist die Implikation

„Wenn $1 + 1 = 2$ ist, dann beträgt die Summe der Innenwinkel aller Dreiecke 180° .“

wahr, obwohl die Gleichung $1 + 1 = 2$ wenig bis nichts mit den geometrischen Eigenschaften irgendwelcher Dreiecke zu tun hat. Ausschlaggebend für den Wahrheitsgehalt der Implikation ist hier nur, dass die beiden Teilaussagen wahr sind.

Implikationen spielen in der Mathematik eine sehr wichtige Rolle; so gut wie jeder mathematische Satz wird als Implikation formuliert. Im Mathematikunterricht werden Implikationen bereits bei ganz elementaren Vorgängen wie etwa der **Umformung** von Gleichungen verwendet. So verwendet man beispielsweise die Tatsache, dass die Implikation „ $x + 3 = 5 \Rightarrow x = 2$ “ für alle reellen Zahlen x gültig ist, um die Gleichung $x + 3 = 5$ nach x hin „aufzulösen“.

- (v) **Äquivalenz** $\varphi \Leftrightarrow \psi$ „Es gilt φ genau dann, wenn ψ gilt.“ / „ φ ist hinreichende und zugleich notwendige Bedingung für ψ .“

φ	ψ	$\varphi \Leftrightarrow \psi$
w	w	w
w	f	f
f	w	f
f	f	w

Beim Arbeiten mit Implikationen ist es sehr wichtig, die zusammengesetzten Aussagen „ $\varphi \Rightarrow \psi$ “, „ $\psi \Rightarrow \varphi$ “ und „ $\varphi \Leftrightarrow \psi$ “ sorgfältig auseinander zu halten. Geschieht dies nicht, dann kann das bereits beim Auflösen von quadratischen Gleichungen zu Fehlern führen. Beispielsweise ist die Implikation $x = 3 \Rightarrow x^2 = 9$ für alle reellen Zahlen x gültig, während $x^2 = 9 \Rightarrow x = 3$ für alle reellen Zahlen $x \neq -3$ richtig, für $x = -3$ aber falsch ist: In diesem Fall ist Prämisse $x^2 = 9$ wahr, die Konklusion $x = 3$ aber falsch, damit ist die gesamte Implikation falsch. Wendet man nun diese fehlerhafte Implikation bei der Auflösung der Gleichung $x^2 - 8x + 7 = 0$ an, so erhält man

$$x^2 - 8x + 7 = 0 \Rightarrow x^2 - 8x = -7 \Rightarrow x^2 - 8x + 16 = 9 \Rightarrow (x - 4)^2 = 9$$

$$\Rightarrow x - 4 = 3 \Rightarrow x = 7$$

und „verliert“ somit die Lösung $x = 1$ der Gleichung. Der Fehler tritt an der Stelle auf, wo das Implikationszeichen $\Rightarrow$ in Anführungsstriche gesetzt wurde. Man könnte auch sagen, dass der Fehler in der Rechnung oben dadurch zu Stande kam, dass an einer Stelle eine notwendige Bedingung mit einer hinreichenden Bedingung verwechselt wurde: Die Gleichung $(x - 4)^2 = 9$ ist zwar eine notwendige Bedingung dafür, dass $x - 4 = 3$ ist, aber eben keine hinreichende. Dieser Unterschied spielt, wie wir noch sehen werden, bei vielen mathematischen Sätzen eine wichtige Rolle, zum Beispiel bei der Bestimmung von lokalen Extremstellen einer Funktion.

Häufig werden durch **mehrfache** Anwendung der Verknüpfungssymbole nicht nur zwei, sondern mehrere Aussagen miteinander verbunden. In welcher Reihenfolge dies geschieht, wird durch Klammern festgelegt. Beispielsweise bedeutet $(\varphi \wedge \psi) \vee \rho$, dass zuerst φ und ψ miteinander „und“-verknüpft und diese Aussage dann anschließend mit der Aussage ρ noch „oder“-verknüpft wird.

Um Schreibarbeit (also Klammern) einzusparen, legt man fest, dass bestimmte Symbole stärker binden als andere, vergleichbar mit der Konvention „Punktrechnung vor Strichrechnung“ aus der Arithmetik. Per Festlegung bindet das Negationszeichen $\neg$ am stärksten, danach in absteigender Reihenfolge die Zeichen $\wedge$, $\vee$, $\Rightarrow$ und $\Leftrightarrow$. Beispielsweise ist der Ausdruck

$$\neg\varphi \wedge \neg\psi \Rightarrow \varphi \Leftrightarrow \psi \quad \text{gleichbedeutend mit} \quad (((\neg\varphi) \wedge (\neg\psi)) \Rightarrow \varphi) \Leftrightarrow \psi.$$

Gelegentlich kann der Wahrheitsgehalt einer zusammengesetzten Aussage bestimmt werden, ohne dass man die Teilaussagen, aus denen die Aussage besteht, überhaupt kennt. Solche Aussagen wirken auf den ersten Blick eher nutzlos, bilden aber die Grundlage für das **logische Schließen** innerhalb einer mathematischen Beweisführung.

(1.2) Definition Eine zusammengesetzte Aussage, die unabhängig vom Wahrheitsgehalt ihrer Teilaussagen immer wahr ist, wird **Tautologie** genannt.

Ein Beispiel für eine Tautologie ist die bekannte Bauernregel

„Wenn der Hahn kräht auf dem Mist, dann ändert sich das Wetter, oder es bleibt, wie es ist.“

Isolieren wir hier die Teilaussagen

φ = „Der Hahn kräht auf dem Mist.“
 ψ = „Das Wetter ändert sich.“

und interpretieren den Satz „Das Wetter bleibt, wie es ist.“ als Negation $\neg\psi$ von ψ , dann ist unsere Bauernregel ϕ in Kurzschreibweise durch $\varphi \Rightarrow (\psi \vee \neg\psi)$ gegeben. Wir wissen bereits, dass der Wahrheitsgehalt von ϕ nur von den Wahrheitswerten der Aussagen φ und ψ abhängt. Um zu kontrollieren, ob es sich bei ϕ um eine Tautologie handelt, genügt es also, alle möglichen Kombinationen von Wahrheitswerten für φ und ψ in den Ausdruck ϕ einzusetzen. Wir erledigen dies durch Ausfüllen einer Tabelle.

φ	ψ	$\neg\psi$	$\psi \vee \neg\psi$	ϕ
w	w	f	w	w
w	f	w	w	w
f	w	f	w	w
f	f	w	w	w

Die zusammengesetzte Aussage ϕ ist unabhängig von φ und ψ immer wahr, also eine Tautologie.

(1.3) Definition Wir sagen, die Aussage ψ folgt aus den Aussagen $\varphi_1, \dots, \varphi_n$ durch einen **logischen Schluss**, wenn die Implikation

$\varphi_1 \wedge \dots \wedge \varphi_n \Rightarrow \psi$ eine Tautologie ist.

Wir sehen uns nun eine Reihe von logischen Schlüssen an, die in der Mathematik häufig verwendet werden. Im folgenden bezeichnen φ , ϕ und ψ jeweils beliebige Aussagen. Mit Hilfe von Wahrheitstabellen wird überprüft, dass die logischen Schlüsse zulässig sind.

(i) **Modus Ponens** $\varphi \wedge (\varphi \Rightarrow \psi) \Rightarrow \psi$

„Wenn φ gilt und aus φ die Aussage ψ folgt, dann gilt ψ .“

φ	ψ	$\varphi \Rightarrow \psi$	$\varphi \wedge (\varphi \Rightarrow \psi)$	$\varphi \wedge (\varphi \Rightarrow \psi) \Rightarrow \psi$
w	w	w	w	w
f	w	w	f	w
w	f	f	f	w
f	f	w	f	w

(ii) **Beweis durch Kontraposition** $(\varphi \Rightarrow \psi) \Leftrightarrow (\neg\psi \Rightarrow \neg\varphi)$

„Aus φ folgt ψ genau dann, wenn aus $\neg\psi$ die Aussage $\neg\varphi$ folgt.“

φ	ψ	$\neg\varphi$	$\neg\psi$	$\varphi \Rightarrow \psi$	$\neg\psi \Rightarrow \neg\varphi$	$(\varphi \Rightarrow \psi) \Leftrightarrow (\neg\psi \Rightarrow \neg\varphi)$
w	w	f	f	w	w	w
w	f	f	w	f	f	w
f	w	w	f	w	w	w
f	f	w	w	w	w	w

(iii) **Beweis durch Widerspruch** $(\neg\varphi \Rightarrow \phi \wedge \neg\phi) \Rightarrow \varphi$

„Wenn aus $\neg\varphi$ ein Widerspruch folgt (nämlich eine Aussage ϕ und zugleich auch ihr Gegenteil $\neg\phi$), dann ist φ wahr.“

φ	ϕ	$\neg\varphi$	$\neg\phi$	$\phi \wedge \neg\phi$	$\neg\varphi \Rightarrow \phi \wedge \neg\phi$	$(\neg\varphi \Rightarrow \phi \wedge \neg\phi) \Rightarrow \varphi$
w	w	f	f	f	w	w
w	f	f	w	f	w	w
f	w	w	f	f	f	w
f	f	w	w	f	f	w

(iv) **Satz vom Ringschluss** $(\varphi \Rightarrow \phi) \wedge (\phi \Rightarrow \psi) \wedge (\psi \Rightarrow \varphi) \Rightarrow (\varphi \Leftrightarrow \phi) \wedge (\phi \Leftrightarrow \psi) \wedge (\psi \Leftrightarrow \varphi)$

„Wenn aus φ die Aussage ϕ und aus ϕ die Aussage ψ und aus ψ wieder die Aussage φ folgt, dann sind die drei Aussagen φ , ϕ und ψ äquivalent.“

Hier ist die Verifikation etwas aufwändiger als bei den vorherigen Regeln. Zur Abkürzung definieren wir die Teilaussagen $A = \varphi \Rightarrow \phi$, $B = \phi \Rightarrow \psi$, $C = \psi \Rightarrow \varphi$, $D = \varphi \Leftrightarrow \phi$, $E = \phi \Leftrightarrow \psi$ und $F = \psi \Leftrightarrow \varphi$.

Damit erhalten wir

φ	ϕ	ψ	A	B	C	D	E	F	$A \wedge B \wedge C$	$D \wedge E \wedge F$	$A \wedge B \wedge C \Rightarrow D \wedge E \wedge F$
w	w	w	w	w	w	w	w	w	w	w	w
w	w	f	w	f	w	w	f	f	f	f	w
w	f	w	f	w	w	f	f	w	f	f	w
w	f	f	f	w	w	f	w	f	f	f	w
f	w	w	w	w	f	f	w	f	f	f	w
f	w	f	w	f	w	f	f	w	f	f	w
f	f	w	w	w	f	w	f	f	f	f	w
f	f	f	w	w	w	w	w	w	w	w	w

§ 2. Mengenlehre und Prädikatenlogik

Inhaltsübersicht

Fast die gesamte moderne Mathematik ist auf dem Begriff der **Menge** aufgebaut. Eine Menge kann durch Aufzählung ihrer Elemente oder durch eine definierende Bedingung, ein sog. **Aussagenschema**, beschrieben werden. Mit Hilfe von Aussagenschemata definieren wir auch einige wichtige **Mengenoperationen**. Außerdem werden mit ihnen **quantifizierte** Aussagen gebildet, wie sie bei der Formulierung mathematischer Sätze fast immer vorkommen. Zum Abschluss führen wir die natürlichen Zahlen ein und besprechen das Prinzip der **vollständigen Induktion**.

Wichtige Begriffe und Sätze

- Mengendefinition nach Cantor
- Bedeutung der Relationen $\in$, $\subseteq$, $\supseteq$, $\subsetneq$, $\supsetneq$
- Definition von Mengen durch definierende Bedingungen (Aussagenschemata)
- Mengenoperationen (Durchschnitt, Vereinigung, Differenz, kartesisches Produkt, Potenzmengenbildung)
- Nachweis der Mengengleichheit
- quantifizierte Aussagen, All- und Existenzquantor ($\forall$, $\exists$)
- Prinzip der vollständigen Induktion

Fast die gesamte moderne Mathematik basiert auf dem Konzept der Menge. Dies bedeutet, dass fast jedes mathematische Objekt, egal ob es sich dabei um eine Zahl, eine Funktion oder ein geometrisches Gebilde handelt, letztendlich durch eine Menge beschrieben werden kann. Desweiteren kann fast jede mathematische Aussage auf die Mengenlehre zurückgeführt und mit den Mitteln der Mengenlehre bewiesen werden, eine ganz erstaunliche Feststellung, wenn man sich die Vielfalt und Verschiedenartigkeit der mathematischen Strukturen vor Augen hält.

(2.1) Definition (naive Mengendefinition von Cantor)

„Eine **Menge** ist eine beliebige Zusammenfassung von bestimmten, wohlunterschiedenen Objekten unserer Anschauung oder unseres Denkens – welche die **Elemente** dieser Menge genannt werden – zu einem Ganzen.“

Hierbei handelt es sich nicht um eine Definition im streng mathematischen Sinn; Begriffe wie „Zusammenfassung“, „Objekt“, „Anschauung“ usw. werden ihrerseits nicht definiert, sondern rein intuitiv verwendet. Auf Grund unserer Alltagserfahrung ist die Bedeutung der Cantorschen Definition dennoch unmittelbar klar. Jeder kann sich vorstellen, was es heißt, „Objekte unserer Anschauung“ zu einem „Ganzen“ zusammenzufassen (z.B. die Bürger einer Gemeinde, die Möbelstücke in einer Wohnung, die Moleküle eines Wassertropfens usw.), dasselbe gilt für die „Objekte unseres Denkens“ wie etwa die natürlichen Zahlen oder geometrische Figuren.

Wir weisen auf zwei wichtige Punkte der Cantorschen Definition hin: Erstens sind sämtliche Objekte einer Menge **verschieden**, es ist also nicht möglich, dass ein und dasselbe Objekt mehrfach in einer Menge vorkommt. Zweitens ist jede Menge als „Zusammenfassung“ durch ihre Elemente **eindeutig bestimmt**. Dies bedeutet, dass zwei Mengen genau dann gleich sind, wenn sie dieselben Elemente enthalten.

Folgende Kurzschreibweisen sind in der Mengenlehre üblich.

$x \in M$	Das Objekt x ist Element der Menge M .
$x \notin M$	Das Objekt x ist <i>kein</i> Element der Menge M , in Kurzform also $\neg(x \in M)$.
$M \subseteq N$	Jedes Element x von M ist auch ein Element von N , d.h. die Implikation $x \in M \Rightarrow x \in N$ ist für alle Objekte x erfüllt. Man bezeichnet M dann als Teilmenge von N .
$M = N$	Es gilt $x \in M \Leftrightarrow x \in N$ für alle Objekte x (äquivalent: $M \subseteq N \wedge N \subseteq M$).
$M \supseteq N$	gleichbedeutend mit $N \subseteq M$
$M \subsetneq N$	$M \subseteq N \wedge \neg(M = N)$
$M \supsetneq N$	gleichbedeutend mit $N \subsetneq M$
$\emptyset$	die leere Menge Dies ist die eindeutig bestimmte Menge die $x \notin \emptyset$ für alle Objekte x erfüllt, also die Menge, die kein einziges Objekt als Element besitzt.

Es gibt mehrere Möglichkeiten, eine Menge konkret anzugeben. Zunächst kann dies **umgangssprachlich** geschehen.

„Sei P die Menge aller Primzahlen.“

Eine andere Möglichkeit besteht darin, die Elemente einer Menge explizit **aufzuzählen**.

$$M = \{1, 2, 3, 4, 5, 6, 7\} \quad \text{oder kürzer} \quad M = \{1, 2, \dots, 7\}$$

Bei der Verwendung von „...“ ist darauf zu achten, dass für den Leser klar ersichtlich ist, welche Elemente bei der Aufzählung weggelassen wurden. Schreibt man etwa $P = \{2, 3, 5, 7, 11, 13, 17, \dots\}$, dann ist noch einigermaßen ersichtlich, dass die Menge der Primzahlen gemeint ist. Schwieriger wird das schon bei der Angabe

$$M = \{1, 4, \dots, 64\}.$$

Hier ist nicht ohne weiteres klar, ob die Menge $\{1, 4, 16, 64\} = \{4^0, 4^1, 4^2, 4^3\}$ der ersten drei Viererpotenzen oder vielleicht $\{1, 4, 9, 16, 25, 36, 49, 64\} = \{1^2, 2^2, 3^2, 4^2, 5^2, 6^2, 7^2, 8^2\}$, die Menge der ersten acht Quadratzahlen, gemeint ist. Ein Vorteil der „...“-Schreibweise besteht aber darin, dass mit ihr auch unendliche Mengen direkt angegeben werden können, zum Beispiel die Menge $\mathbb{N} = \{1, 2, 3, \dots\}$ der natürlichen Zahlen.

Auch mit Hilfe der im letzten Kapitel eingeführten Aussagenschemata lassen sich Mengen definieren. Sei φ ein Aussagenschema mit einem Parameter x und M eine Menge. Für jedes $c \in M$ bezeichnen wir mit $\varphi(c)$ den Satz, den man erhält, wenn der Parameter x durch c ersetzt wird. Wir setzen voraus, dass $\varphi(c)$ für jedes $c \in M$ eine sinnvolle Aussage ist. Nach Definition besteht dann die Menge

$$N = \{c \in M \mid \varphi(c)\}$$

aus genau denjenigen Elementen c von M , für die $\varphi(c)$ eine wahre Aussage ist. Man nennt φ dann auch die **definierende Bedingung** für die Teilmenge N von M .

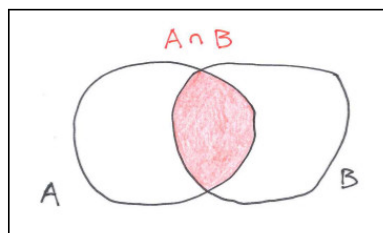
Beispielsweise beschreibt $\{c \in \mathbb{R} \mid c^2 < 1\}$ die Menge derjenigen reellen Zahlen, deren Quadrat kleiner als 1 ist. In dieser Situation ist also $M = \mathbb{R}$ die Grundmenge und $\varphi(x) = x^2 < 1$ das Aussagenschema, dass die Teilmenge beschreibt. Offenbar ist $c^2 < 1$ für jedes $c \in \mathbb{R}$ eine sinnvolle, aber nur für $-1 < c < 1$ auch eine wahre Aussage.

Gelegentlich verwendet man auch die Notation $N = \{c \mid \varphi(c)\}$, ohne die Angabe einer Grundmenge für die Objekte c . In diesem Fall besteht N aus *allen* mathematischen Objekten c , für die die Aussage $\varphi(c)$ wahr ist. Strenggenommen ist eine solche Definition für beliebige Aussagenschemata φ nicht zulässig, weil dies zu Widersprüchen führen kann (Stichwort „Russelsche Antinomie“). Wir werden die Notation aber nur in Situationen einsetzen, wo solche Probleme ausgeschlossen sind.

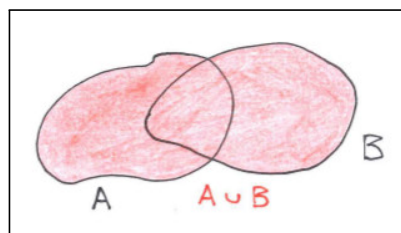
Aus gegebenen Mengen können durch weitere Operationen neue Mengen definiert werden. Seien A und B beliebig vorgegebene Mengen. Folgende Mengenoperationen sind in der Mathematik allgemein gebräuchlich.

Durchschnitt	$A \cap B = \{a \mid a \in A \wedge a \in B\}$
Vereinigung	$A \cup B = \{a \mid a \in A \vee a \in B\}$
Differenz	$A \setminus B = \{a \mid a \in A \wedge a \notin B\}$
kartesisches Produkt	$A \times B = \{(a, b) \mid a \in A \wedge b \in B\}$
Potenzmenge	$\mathcal{P}(A) = \{B \mid B \subseteq A\}$

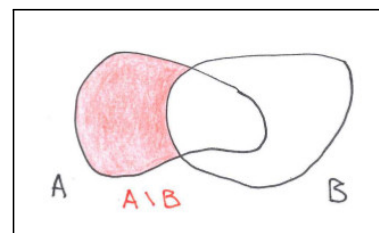
Einige dieser Operationen lassen sich durch sog. **Venn-Diagramme** veranschaulichen.



Durchschnitt



Vereinigung



Differenz

Das kartesische Produkt $A \times B$ besteht aus allen Paaren (a, b) , die mit Elementen $a \in A$ und $b \in B$ gebildet werden können. Ist beispielsweise $A = \{1, 2, 3\}$ und $B = \{1, 2, 4, 5\}$, dann erhalten wir

$$A \times B = \{1, 2, 3\} \times \{1, 2, 4, 5\} = \left\{ \begin{array}{cccc} (1, 1), & (1, 2), & (1, 4), & (1, 5), \\ (2, 1), & (2, 2), & (2, 4), & (2, 5), \\ (3, 1), & (3, 2), & (3, 4), & (3, 5) \end{array} \right\}$$

(Die Elemente wurden nur zur besseren Übersicht in einem rechteckigen Schema angeordnet. Man hätte auch alle 12 Elemente direkt hintereinander schreiben können.)

Bei der Definition des kartesischen Produkts ist zu beachten, dass zwei **Paare** (a, b) und (c, d) von Objekten a, b, c, d nur dann gleich sind, wenn $a = c$ und $b = d$ erfüllt ist. Zum Beispiel sind die Paare $(3, 5)$ und $(5, 3)$ verschieden. Im Gegensatz dazu stimmen die zweielementigen Mengen $\{3, 5\}$ und $\{5, 3\}$ überein, da es keine Rolle spielt, in welcher Reihenfolge die Elemente einer Menge aufgezählt werden.

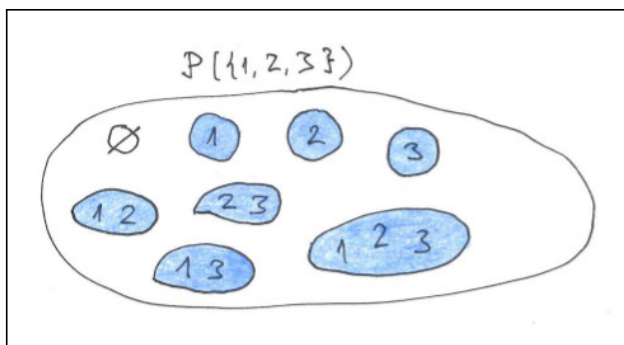
Das kartesische Produkt kann auch mit mehr als zwei Mengen gebildet werden. Die Elemente bezeichnet man dann nicht mehr als Paare, sondern als **Tupel**. Sind beispielsweise A, B, C drei beliebige Mengen, dann setzt man

$$A \times B \times C = \{ (a, b, c) \mid a \in A, b \in B, c \in C \}$$

wobei wieder $(a, b, c) = (a', b', c')$ nur dann erfüllt ist, wenn $a = a'$, $b = b'$ und $c = c'$ gilt. Es ist also beispielsweise $(1, 2, 3) \neq (1, 3, 2)$ und $(2, 2, 4) \neq (2, 4, 2)$. Häufig werden mehrfache kartesische Produkte auch mit ein- und derselben Menge gebildet. Man definiert $A^2 = A \times A$, $A^3 = A \times A \times A$, $A^4 = A \times A \times A \times A$ usw. Beispielsweise ist $(3, 4, \frac{1}{2}, \sqrt{2}, -9)$ ein Element der Menge $\mathbb{R}^5$.

Bei den Potenzmengen ist zu beachten, dass deren Elemente selbst wieder Mengen sind! Nach Definition ist für jede beliebige Mengen A, B die Aussage $B \in \mathcal{P}(A)$ äquivalent zu $B \subseteq A$. Intuitiv klar ist, dass bei einer endlichen Menge A die Potenzmenge $\mathcal{P}(A)$ ebenfalls nur endlich viele Elemente enthält. Ist beispielsweise $A = \{1, 2, 3\}$, dann enthält jede Teilmenge von A entweder kein, genau ein, genau zwei oder genau drei Elemente. Dies kann für eine systematische Aufzählung der Elemente von $\mathcal{P}(A)$ verwendet werden: Es gilt

$$\mathcal{P}(A) = \{ \emptyset, \{1\}, \{2\}, \{3\}, \{1, 2\}, \{1, 3\}, \{2, 3\}, \{1, 2, 3\} \}.$$



die achtelementige Potenzmenge von $\{1, 2, 3\}$

Eine häufige Fehlerquelle beim Umgang mit Mengen besteht darin, dass man zwischen einer Menge und ihren Elementen nicht klar unterscheidet. Beispielsweise wäre es falsch zu sagen, dass die 1 ein Element von $\mathcal{P}(A)$ ist. Lediglich die Menge $\{1\}$ ist ein Element von $\mathcal{P}(A)$, und 1 ist ein Element der Menge $\{1\}$. Momentan klingt das noch recht haarspalterisch; bei komplizierteren mengentheoretischen Konstruktionen (zum Beispiel Faktorstrukturen) kommt man aber in große Schwierigkeiten, wenn man sich an diese Unterscheidung nicht gewöhnt hat.

Eine wichtige Grundtechnik beim Führen von Beweisen ist der Nachweis der **Gleichheit zweier Mengen**. Häufig bietet es sich an, die Aussage $M = N$ in die folgenden beiden Teilaussagen zu zerlegen und diese einzeln zu beweisen.

- (i) Ist x ein Element der Menge M , dann ist x auch ein Element von N .
- (ii) Ist x ein Element der Menge N , dann ist x auch ein Element von M .

Aus der ersten Aussage folgt $M \subseteq N$, aus der zweiten $N \subseteq M$, insgesamt also $M = N$.

Wie die Beweise der Teilaussagen (i) und (ii) aussehen können, schauen wir uns an einem konkreten Beispiel an. Unser Ziel ist der Beweis der Gleichung

$$\{(x, y, z) \in \mathbb{R}^3 \mid (xy + 1)z = 0\} = \{(x, y, 0) \mid x, y \in \mathbb{R}\} \cup \{(x, -\frac{1}{x}, z) \mid x \in \mathbb{R} \setminus \{0\}, z \in \mathbb{R}\}.$$

Wir bezeichnen die Menge auf der linken Seite der Gleichung mit M und die Menge auf der rechten Seite der Gleichung mit N . Die Menge N enthält also alle Tupel der Form $(x, y, 0)$ mit $x, y \in \mathbb{R}$ und alle Tupel der Form $(x, -\frac{1}{x}, z)$ mit $x, z \in \mathbb{R}$ und $x \neq 0$.

Beweis der Teilaussage (i)

Sei $p \in M$. Nach Definition von M gilt $p = (x, y, z) \in \mathbb{R}^3$ mit $x, y, z \in \mathbb{R}$ und $(xy + 1)z = 0$. Aus dieser Gleichung folgt $xy + 1 = 0$ oder $z = 0$, denn das Produkt zweier reeller Zahlen ist nur dann gleich Null, wenn einer der beiden Faktoren gleich Null ist. Ist $z = 0$, dann hat p die Form $(x, y, 0)$ mit $x, y \in \mathbb{R}$, also liegt p in N . Ist dagegen $xy + 1 = 0$, dann muss $x \neq 0$ gelten, denn ansonsten wäre $xy + 1 = 0 \cdot y + 1 = 1$. Aus $xy + 1 = 0$ folgt $xy = -1$, wegen $x \neq 0$ dann auch $y = -\frac{1}{x}$ und somit ebenfalls $p = (x, y, z) = (x, -\frac{1}{x}, z) \in N$.

Beweis der Teilaussage (ii)

Sei $p \in N$. Dann gilt $p = (x, y, 0)$ mit geeigneten $x, y \in \mathbb{R}$ oder $p = (x, -\frac{1}{x}, z)$ für geeignete $x \in \mathbb{R} \setminus \{0\}$ und ein $z \in \mathbb{R}$. Betrachten wir zunächst den Fall $p = (x, y, 0)$ mit $x, y \in \mathbb{R}$. Setzen wir $z = 0$, dann gilt $p = (x, y, z) \in \mathbb{R}^3$ und $(xy + 1)z = (xy + 1) \cdot 0 = 0$. Daraus folgt $p \in M$. Betrachten wir nun den Fall, dass $p = (x, -\frac{1}{x}, z)$ für ein $x \in \mathbb{R} \setminus \{0\}$ und ein $z \in \mathbb{R}$ gilt. Setzen wir $y = -\frac{1}{x}$, dann liegt $p = (x, y, z)$ in $\mathbb{R}^3$. Außerdem gilt $(xy + 1)z = (x \cdot (-\frac{1}{x}) + 1)z = ((-1) + 1)z = 0 \cdot z = 0$. Also liegt p auch in diesem Fall in M . $\square$

In einfacheren Situationen lässt sich die Gleichheit zweier Mengen auch durch eine Kette von Äquivalenzumformungen beweisen. Als Beispiel betrachten wir die Mengengleichung

$$\{x \in \mathbb{R} \mid x^2 + x - 6 = 0\} = \{-3, 2\}.$$

Wieder sei M die Menge auf der linken und N die Menge auf der rechten Seite der Gleichung. Es gilt $M = N$, wenn wir für jedes Objekt x die Äquivalenz $x \in M \Leftrightarrow x \in N$ beweisen können. Da M und N beides Teilmengen von $\mathbb{R}$ sind, genügt es, die Äquivalenz für alle $x \in \mathbb{R}$ zu beweisen, denn ansonsten sind die Aussagen $x \in M$ und $x \in N$ beide falsch und die Äquivalenz damit auf jeden Fall wahr.

Sei also $x \in \mathbb{R}$ vorgegeben. Dann gilt

$$\begin{aligned} x \in M &\Leftrightarrow x^2 + x - 6 = 0 \Leftrightarrow x^2 + x = 6 \Leftrightarrow x^2 + x + \frac{1}{4} = \frac{25}{4} \Leftrightarrow \left(x + \frac{1}{2}\right)^2 = \left(\frac{5}{2}\right)^2 \\ &\Leftrightarrow \left(x + \frac{1}{2}\right)^2 - \left(\frac{5}{2}\right)^2 = 0 \Leftrightarrow \left(\left(x + \frac{1}{2}\right) + \frac{5}{2}\right)\left(\left(x + \frac{1}{2}\right) - \frac{5}{2}\right) = 0 \Leftrightarrow (x + 3)(x - 2) = 0 \\ &\Leftrightarrow x + 3 = 0 \vee x - 2 = 0 \Leftrightarrow x = -3 \vee x = 2 \Leftrightarrow x \in \{-3, 2\} \Leftrightarrow x \in N. \end{aligned}$$

Hier wurde nichts anderes getan, als die Gleichung durch Bildung der quadratischen Ergänzung zu lösen. Wichtig ist bei solchen Beweisen, dass jeder einzelne Schritt genau begründet werden kann, und dass jeweils *beide* Implikationsrichtungen gültig sind. Beispielsweise wäre $x = 3 \Leftrightarrow x^2 = 9$ keine zulässige Äquivalenzumformung, weil die Implikationsrichtung „ $\Leftarrow$ “ nicht für jede reelle Zahl x gültig ist. (Wie wir bereits oben festgestellt haben, ist sie für $x = -3$ falsch.)

In vielen Situationen möchte man Aussagen formulieren, die die Gesamtheit der Elemente einer Menge betreffen. Dazu verwendet man den sog. **Allquantor** $\forall$ und den **Existenzquantor** $\exists$. Sei φ ein Aussagenschema mit x als Parameter, und wiederum sei M eine Menge mit der Eigenschaft, dass man für jedes $c \in M$ durch Ersetzung von x durch c eine sinnvolle Aussage $\varphi(c)$ erhält. Dann kann man mit Hilfe von All- und Existenzquantor zwei neue Aussagen $\forall x \in M : \varphi$ und $\exists x \in M : \varphi$ bilden, die man als **quantifizierte** Aussagen bezeichnet. Den Umgang mit quantifizierten Aussagen bezeichnet man als **Prädikatenlogik**, im Unterschied zur Aussagenlogik, wo man nur Aussagen ohne Quantoren betrachtet.

(2.2) Definition

- (i) Die Aussage $\forall x \in M : \varphi$ bedeutet, dass $\varphi(c)$ für **alle** $c \in M$ wahr ist.
Es gilt also $\{c \in M \mid \varphi(c)\} = M$.
- (ii) Die Aussage $\exists x \in M : \varphi$ bedeutet, dass $\varphi(c)$ für **mindestens ein** $c \in M$ wahr ist.
Es gilt also $\{c \in M \mid \varphi(c)\} \neq \emptyset$.

Betrachten wir beispielsweise das Aussagenschema $x \leq 5$ mit dem Parameter x über der Menge $M = \mathbb{R}$ der reellen Zahlen und bezeichnen es mit φ .

- (i) Die Aussage $\forall x \in \mathbb{R} : x \leq 5$ ist **falsch**, denn $\varphi(c)$ ist nicht für alle $c \in \mathbb{R}$ erfüllt. Beispielsweise ist $\varphi(7)$ falsch.
- (ii) Die Aussage $\exists x \in \mathbb{R} : x \leq 5$ ist **wahr**, denn es gibt Elemente $c \in \mathbb{R}$, für die $\varphi(c)$ wahr ist. Zum Beispiel ist $\varphi(4)$ eine wahre Aussage.

Die meisten Aussagen, die wir im Laufe der Zeit beweisen werden, sind quantifizierte Aussagen, enthalten also die Formulierungen „für alle“ oder „es gibt ein x , so dass...“. Dabei treten besonders zu Anfang häufig methodische Fehler auf. Um eine Aussage der Form $\forall x \in M : \varphi(x)$ zu beweisen, muss die Aussage $\varphi(c)$ für **jedes** $c \in M$ bewiesen werden. Dazu gibt man sich mit der Floskel „Sei $c \in M$.“ ein beliebiges Element c aus M vor und beweist anschließend die Aussage $\varphi(c)$. Während des Beweises darf dann nur verwendet werden, dass c ein Element der Menge M ist. Jede Einschränkung oder Spezialisierung von c macht den Beweis **ungültig**.

Um andererseits eine Aussage der Form $\exists x \in M : \varphi(x)$ zu beweisen, genügt es, ein **spezielles** Element $c \in M$ anzugeben und die Aussage $\varphi(c)$ nur für dieses c zu beweisen. Natürlich kann es schwierig sein, ein solches c erst einmal zu finden. Um beispielsweise die Aussage $\exists x \in \mathbb{R} : x^2 + x - 6 = 0$ auf diesem Weg zu beweisen, muss eine Lösung der quadratischen Gleichung $x^2 + x - 6 = 0$ gefunden werden.

Gelegentlich hat man es auch mit der **Negation** einer quantifizierten Aussage zu tun.

(2.3) Satz Sei M eine Menge und $\varphi(x)$ ein Aussagenschema mit der Eigenschaft, dass $\varphi(c)$ für jedes $c \in M$ eine sinnvolle Aussage ist. Dann gelten die folgenden Äquivalenzen.

$$(i) \quad \neg \forall x \in M : \varphi(x) \iff \exists x \in M : \neg \varphi(x)$$

$$(ii) \quad \neg \exists x \in M : \varphi(x) \iff \forall x \in M : \neg \varphi(x)$$

Beweis: zu (i) „ $\Leftarrow$ “ (durch Widerspruch) Auf Grund der Voraussetzung gibt es ein $c \in M$, so dass $\neg \varphi(c)$ erfüllt ist. Nehmen wir nun an, auch die Aussage $\forall x \in M : \varphi(x)$ ist wahr. Dann muss insbesondere auch $\varphi(c)$ gelten. Die Aussagen $\varphi(c)$ und $\neg \varphi(c)$ wären also gleichzeitig erfüllt, was unmöglich ist. Der Widerspruch zeigt, dass $\neg \forall x \in M : \varphi(x)$ gelten muss.

„ $\Rightarrow$ “ (durch Kontraposition und Widerspruch) Wir setzen $\neg \exists x \in M : \neg \varphi(x)$ voraus und zeigen $\forall x \in M : \varphi(x)$. Sei dazu $c \in M$ vorgegeben. Angenommen, es gilt $\neg \varphi(c)$. Dann existiert also ein $c \in M$ mit $\neg \varphi(c)$, und das bedeutet, dass $\exists x \in M : \neg \varphi(x)$ gilt, im Widerspruch zu unserer Voraussetzung. Da also $\neg \varphi(c)$ zu einem Widerspruch führt, muss $\varphi(c)$ gelten. Weil $c \in M$ beliebig vorgegeben war, haben wir somit $\forall x \in M : \varphi(x)$ bewiesen.

zu (ii) „ $\Leftarrow$ “ (durch Kontraposition und Widerspruch) Wir setzen $\exists x \in M : \varphi(x)$ voraus und leiten daraus $\neg \forall x \in M : \neg \varphi(x)$ ab. Auf Grund unserer Voraussetzung existiert ein $c \in M$, so dass $\varphi(c)$ erfüllt ist. Nehmen wir nun an, es gilt $\forall x \in M : \neg \varphi(x)$. Dann müsste auch $\neg \varphi(c)$ gelten, im Widerspruch zu $\varphi(c)$. Also gilt statt dessen $\neg \forall x \in M : \neg \varphi(x)$.

„ $\Rightarrow$ “ (durch Widerspruch) Unsere Voraussetzung lautet $\neg \exists x \in M : \varphi(x)$. Zum Beweis von $\forall x \in M : \neg \varphi(x)$ sei $c \in M$ vorgegeben. Nehmen wir an, es gilt $\varphi(c)$. Dann wäre die Aussage $\exists x \in M : \varphi(x)$ erfüllt, im Widerspruch zu unserer Voraussetzung. Also gilt $\neg \varphi(c)$. Weil $c \in M$ beliebig vorgegeben war, ist damit $\forall x \in M : \neg \varphi(x)$ bewiesen. $\square$

Schließlich ist es noch möglich, mehrere Quantoren zu **verschachteln**. In diesem Fall benötigt man Aussagenschemata mit **mehreren** Parametern. Sei φ ein Aussagenschema mit den beiden Parametern x, y und M, N Mengen mit der Eigenschaft, dass man für alle $c \in M$ und $d \in N$ eine sinnvolle Aussage $\varphi(c, d)$ erhält, wenn man x durch c und y durch d ersetzt. Dann ist $\forall y \in N : \varphi$ ein Aussagenschema, das nur noch vom Parameter x abhängt, und

$$\exists x \in M : \forall y \in N : \varphi$$

ist eine Aussage, mit zwei ineinander verschachtelten Quantoren. Umgangssprachlich bedeutet diese Aussage: „Es gibt ein $c \in M$, so dass für alle $d \in N$ jeweils $\varphi(c, d)$ gilt.“

Dabei sind die Quantoren $\exists$ und $\forall$ in beliebiger Kombination zugelassen. Es ergeben sich dadurch Aussagen mit unterschiedlicher Bedeutung.

- Der Ausdruck $\forall x \in M : \exists y \in N : \varphi$ bedeutet:
„Für jedes $c \in M$ gibt es ein $d \in N$, so dass $\varphi(c, d)$ gilt.“
- Der Ausdruck $\exists x \in M : \exists y \in N : \varphi$ bedeutet:
„Es gibt Elemente $c \in M$ und $d \in N$, so dass $\varphi(c, d)$ gilt.“
- Der Ausdruck $\forall x \in M : \forall y \in N : \varphi$ bedeutet:
„Für alle $c \in M$ und alle $d \in N$ gilt $\varphi(c, d)$.“

Man beachte, dass auch die Aussagen $\forall x \in M : \exists y \in N : \varphi$ und $\exists y \in N : \forall x \in M : \varphi$ nicht etwa gleichbedeutend sind, es kommt also auch auf die **Reihenfolge** der Quantoren an. Wir machen uns dies am Beispiel des Aussagenschemas $x < y$ mit den Parametern x, y klar, das wir wieder mit φ bezeichnen. Offenbar erhält man jedes Mal eine sinnvolle Aussage, wenn man für x und y Elemente aus $\mathbb{R}$, der Menge der reellen Zahlen, einsetzt.

Die Aussage $\exists x \in \mathbb{R} : \forall y \in \mathbb{R} : x < y$ bedeutet nun: „Es gibt ein $c \in \mathbb{R}$, so dass für alle $d \in \mathbb{R}$ jeweils $c < d$ gilt.“ Diese Aussage ist offenbar falsch. Denn nehmen wir an, es gibt ein solches c . Dann ist $d = c - 1$ offenbar kleiner als c , und nicht größer. Somit ist $c < d$ nicht für alle $d \in \mathbb{R}$ erfüllt.

Die Aussage $\forall y \in \mathbb{R} : \exists x \in \mathbb{R} : x < y$ bedeutet andererseits, dass für jedes $d \in \mathbb{R}$ jeweils ein $c \in \mathbb{R}$ mit $c < d$ existiert. Diese Aussage ist wahr. Geben wir nämlich ein beliebiges $d \in \mathbb{R}$ vor, dann können wir $c = d - 1$ setzen, und die Aussage $c < d$ ist offenbar erfüllt.

Kommen wir nun zum letzten Thema dieses Kapitels, der **vollständigen Induktion**. Wir werden die Menge $\mathbb{N} = \{1, 2, 3, \dots\}$ der natürlichen Zahlen in einem späteren Kapitel als Teilmenge der reellen Zahlen definieren. Trotzdem soll an dieser Stelle bereits mit $\mathbb{N}$ gearbeitet werden. Wir setzen folgende Aussagen über die Menge $\mathbb{N}$ als bekannt voraus. Sie sind in der Literatur unter dem Namen **Peano-Axiome** bekannt, benannt nach dem italienischen Mathematiker *Giuseppe Peano* (1858-1932) und lauten

- (P1) Es gibt ein ausgezeichnetes Element in $\mathbb{N}$, das wir mit 1 bezeichnen.
- (P2) Jedes $n \in \mathbb{N}$ besitzt einen eindeutig bestimmten **Nachfolger**, der mit $n + 1$ bezeichnet wird.
- (P3) Kein Element aus $\mathbb{N}$ besitzt die 1 als Nachfolger.
- (P4) Sind $m, n \in \mathbb{N}$ mit $m + 1 = n + 1$, dann folgt $m = n$.

Hinzu kommt noch das wichtige

(P5) *Induktionsprinzip*:

Sei φ ein Aussagenschema mit folgenden Eigenschaften: Für jedes $n \in \mathbb{N}$ ist $\varphi(n)$ eine sinnvolle Aussage, darüber hinaus seien $\varphi(1)$ und $\forall x \in \mathbb{N} : \varphi(x) \Rightarrow \varphi(x+1)$ wahre Aussagen. Dann ist auch $\forall x \in \mathbb{N} : \varphi(x)$ wahr.

Die Anwendung des Induktionsprinzips bezeichnet man als **vollständige Induktion**, es ist eines der wichtigsten Beweisprinzipien der Mathematik. Wir werden sehen, dass in vielen Situationen die Beweise der Aussagen $\varphi(1)$ und $\forall x \in \mathbb{N} : \varphi(x) \Rightarrow \varphi(x+1)$ erheblich einfacher sind als ein direkter Beweis von $\forall x \in \mathbb{N} : \varphi(x)$.

Wie wir im nächsten Abschnitt sehen werden, kann das Induktionsprinzip zum Beispiel dafür benutzt werden, um auf den natürlichen Zahlen die Addition und die Multiplikation zu definieren, und um die aus der Schule bekannten Rechengesetze herzuleiten, zum Beispiel Assoziativ-, Kommutativ- und Distributivgesetz. Wir nehmen aber hier an, dass wir die aus der Schule bekannten Zahlbereiche schon zur Verfügung haben und betrachten als Beispiel den folgenden „Klassiker“ unter den Induktionsbeweisen.

(2.4) Satz Sei $n \in \mathbb{N}$. Dann ist die Summe der ersten n natürlichen Zahlen gegeben durch

$$1 + 2 + \dots + n = \frac{1}{2}n(n+1).$$

Beweis: Unser Ziel besteht darin, das Induktionsprinzip auf das Aussagenschema $1 + 2 + \dots + x = \frac{1}{2}x(x+1)$ mit dem Parameter x anzuwenden. Bezeichnen wir dieses Aussagenschema mit φ , dann müssen wir also $\varphi(1)$ und $\forall x \in \mathbb{N} : \varphi(x) \Rightarrow \varphi(x+1)$ beweisen. Die Aussage $\varphi(1)$ lautet $1 = \frac{1}{2} \cdot 1 \cdot (1+1)$. Diese Aussage ist offensichtlich wahr, denn tatsächlich gilt $\frac{1}{2} \cdot 1 \cdot (1+1) = \frac{1}{2} \cdot 1 \cdot 2 = 1$.

Nun beweisen wir $\forall x \in \mathbb{N} : \varphi(x) \Rightarrow \varphi(x+1)$. Sei dazu $n \in \mathbb{N}$ vorgegeben. Dann ist $\varphi(n) \Rightarrow \varphi(n+1)$ zu zeigen. Setzen wir dazu voraus, dass $\varphi(n)$ wahr ist. Dann gilt $1 + 2 + \dots + n = \frac{1}{2}n(n+1)$. Diese Gleichung bleibt erhalten, wenn wir auf beiden Seiten $n+1$ addieren, es gilt also $1 + 2 + \dots + n + (n+1) = \frac{1}{2}n(n+1) + (n+1)$. Wegen

$$\begin{aligned} \frac{1}{2}n(n+1) + (n+1) &= \frac{1}{2}n(n+1) + 1 \cdot (n+1) = \left(\frac{1}{2}n+1\right)(n+1) \\ &= \frac{1}{2}(n+2)(n+1) = \frac{1}{2}(n+1)((n+1)+1) \end{aligned}$$

erhalten wir $1 + 2 + \dots + n + (n+1) = \frac{1}{2}(n+1)((n+1)+1)$. Also gilt auch $\varphi(n+1)$. Damit ist insgesamt die Implikation $\varphi(n) \Rightarrow \varphi(n+1)$ bewiesen. $\square$

In einem Induktionsbeweis über ein Aussagenschema φ bezeichnet man den Beweis von $\varphi(1)$ als **Induktionsanfang** und den Beweis der Implikation $\varphi(n) \Rightarrow \varphi(n+1)$ für ein beliebig vorgegebenes $n \in \mathbb{N}$ als **Induktionsschritt**. Die Teilaussage $\varphi(n)$ nennt man dabei die **Induktionsvoraussetzung**. Jeder Induktionsbeweis sollte so aufgeschrieben sein, dass klar zu erkennen ist, an welcher Stelle die Induktionsvoraussetzung verwendet wird. (In einem späteren Kapitel werden wir sehen, wie sich der Ausdruck $1 + 2 + \dots + n$ mit Hilfe des Summenzeichens kompakter darstellen lässt.)

§ 3. Relationen

Inhaltsübersicht

Eine *Relation* auf einer Menge X ist eine Teilmenge des kartesischen Produkts $X \times X$; intuitiv kann man sich darunter eine Beziehung zwischen den Elementen der Menge vorstellen. Wir betrachten in diesem Kapitel zwei wichtige Klassen von Relationen, die *Halbordnungen* und die *Äquivalenzrelationen*. Eine Halbordnung auf X verwendet man, um die Elemente der Menge X auf irgendeine Weise zu „vergleichen“. Von diesem Konzept werden wir vor allem in der Analysis Gebrauch machen. Dagegen dient eine Äquivalenzrelation auf X dazu, die Menge X geeignet zu zerlegen. Dies führt in erster Linie zu Anwendungen in der Algebra, auch in der Linearen Algebra.

Wichtige Begriffe und Sätze

- Relation auf einer Menge X , zwischen zwei Mengen X und Y
- Halbordnungen, Verbände und Totalordnungen
(Beispiel: die Potenzmenge einer Menge als Verband)
- minimales und maximales Element, Minimum und Maximum, Infimum und Supremum
(Beispiel: die Grenzen eines endlichen Intervalls als Infimum und Supremum)
- Äquivalenzrelation auf einer Menge, Äquivalenzklasse
(Beispiel: die Kongruenzrelationen auf $\mathbb{Z}$)
- Korrespondenz zwischen Äquivalenzrelationen und Zerlegungen
(Beispiel: die Zerlegung von $\mathbb{Z}$ in Kongruenzklassen)

(3.1) Definition Seien X und Y Mengen.

- (i) Eine **Relation** auf X ist eine Teilmenge $R \subseteq X \times X$.
- (ii) Eine Relation zwischen X und Y ist eine Teilmenge $R \subseteq X \times Y$.

Intuitiv kann man sich eine Relation R als *Beziehung* zwischen den Elementen von X und Y vorstellen. Für beliebige $x \in X$ und $y \in Y$ soll $(x, y) \in R$ genau dann gelten, wenn x und y miteinander in Beziehung stehen.

Betrachten wir als erstes Beispiel die Relation $|$ auf der Menge $X = \{1, 2, 3, 4, 5, 6\}$ gegeben durch $| = \{(a, b) \in X \times X \mid a \text{ ist Teiler von } b\}$. Man bezeichnet diese Relation als **Teilerrelation**. In ausgeschriebener Form handelt es sich um die Menge

$$| = \{(1, 1), (1, 2), (1, 3), (1, 4), (1, 5), (1, 6), (2, 2), (2, 4), (2, 6), (3, 3), (3, 6), (4, 4), (5, 5), (6, 6)\} \subseteq X \times X.$$

Alternativ könnte man die Relation $|$ auch in Tabellenform darstellen, wobei man für jedes Element von X eine Zeile und eine Spalte vorsieht und an der Position (x, y) genau dann ein Kreuz $\times$ setzt, wenn $(x, y) \in |$ gilt. Dies würde dann folgendermaßen aussehen.

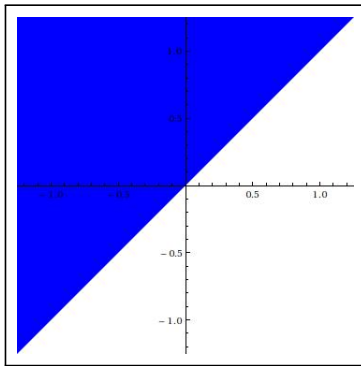
	1	2	3	4	5	6
1	X	X	X	X	X	X
2		X		X		X
3			X			X
4				X		
5					X	
6						X

Natürlich kann man die Teilerrelation auch auf der gesamten Menge $\mathbb{N}$ der natürlichen Zahlen betrachten. Da $\mathbb{N}$ aber unendlich ist, kann man die Elemente von $\mathbb{N}$ natürlich nicht mehr einzeln angeben, weder als Aufzählung noch in Tabellenform.

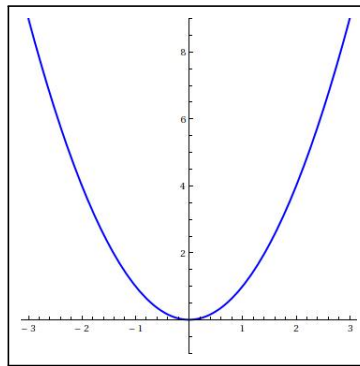
Viele Relationen auf $\mathbb{R}$ lassen sich wiederum graphisch darstellen, weil es sich dabei um nichts anderes als eine Teilmenge der Ebene $\mathbb{R}^2 = \mathbb{R} \times \mathbb{R}$ handelt. So könnte man etwa die Punkte, die zur Relation gehören, in der Ebene blau einzeichnen. Für die Relationen

$$R = \{(x, y) \in \mathbb{R}^2 \mid x \leq y\} \quad , \quad S = \{(x, y) \in \mathbb{R}^2 \mid y = x^2\} \quad \text{und} \quad T = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 \leq 1\}$$

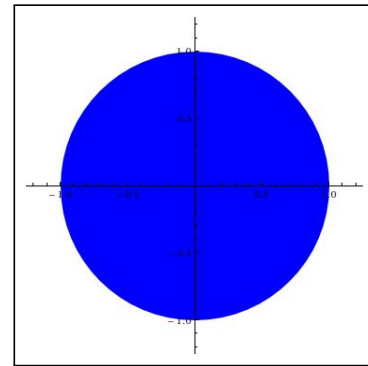
würde man zum Beispiel die folgenden Bilder erhalten.



Relation R



Relation S



Relation T

Beim Umgang mit Relationen verwendet man häufig die folgende **Infix-Notation**: Ist R eine Relation zwischen zwei Mengen X und Y , dann schreibt man die Aussage „ $(x, y) \in R$ “ sehr oft in der Form „ xRy “. Als Bezeichnung für Relationen werden häufig Symbole wie $\leq$, $<$, $\prec$, $\equiv$, $\cong$ usw. verwendet.

Es sei noch darauf hingewiesen, dass eine Relation $R \subseteq X \times Y$ keine bestimmte „Bedeutung“ zu haben braucht, sondern vollkommen willkürlich gewählt werden kann. So ist zum Beispiel auch $R = \{(3, 7), (19, 8), (2, 44)\}$ eine Relation auf $\mathbb{N}$. Um allerdings zu mathematisch „interessanten“ Relationen zu kommen, beschränkt man sich auf Relationen mit bestimmten festgelegten Eigenschaften.

(3.2) Definition Sei X eine Menge. Eine Relation R auf X heißt

- (i) **reflexiv**, falls xRx für alle $x \in R$,
- (ii) **symmetrisch**, falls $xRy \Rightarrow yRx$ für alle $x, y \in R$,
- (iii) **anti-symmetrisch**, falls $xRy \wedge yRx \Rightarrow x = y$ für alle $x, y \in R$, und
- (iv) **transitiv**, falls $xRy \wedge yRz \Rightarrow xRz$ für alle $x, y, z \in R$ gilt.

Wir kommen nun zur Definition der ersten wichtigen großen Klasse von Relationen.

(3.3) Definition Sei X eine Menge.

- (i) Eine **Halbordnung** auf X ist eine reflexive, anti-symmetrische und transitive Relation.
- (ii) Bezeichnet $\leq$ eine Halbordnung auf X , so nennt man zwei Elemente $x, y \in X$ **vergleichbar** bezüglich $\leq$, wenn die Bedingung $(x \leq y) \vee (y \leq x)$ erfüllt ist.
- (iii) Eine Halbordnung auf X wird **Totalordnung** genannt, wenn je zwei Elemente aus X miteinander vergleichbar sind.

Wir betrachten eine Reihe konkreter Beispiele.

- (i) Die gewöhnliche $\leq$ -Relation auf $\mathbb{N}$ ist eine Totalordnung, ebenso die entsprechende Relation auf jeder der Mengen $\mathbb{N}_0$, $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$. Alle Eigenschaften einer Totalordnung (Reflexivität, Anti-Symmetrie, Transitivität, Vergleichbarkeit) lassen sich unmittelbar überprüfen.
- (ii) Die oben beschriebene Teilerrelation $|$ auf der Menge $\mathbb{N}$ ist eine Halbordnung, aber keine Totalordnung.

Zuerst überprüfen wir die Halbordnungseigenschaften. Für jede Zahl $n \in \mathbb{N}$ gilt $n = 1 \cdot n$, also gilt $n | n$. Dies zeigt, dass die Relation reflexiv ist. Zur Überprüfung der Anti-Symmetrie seien $m, n \in \mathbb{N}$ mit $m | n$ und $n | m$ vorgegeben. Nach Definition der Teilbarkeit gibt es $k, \ell \in \mathbb{N}$ mit $n = k \cdot m$ und $m = \ell \cdot n$. Durch Einsetzen erhalten wir $n = k \cdot (\ell \cdot n) = (k \cdot \ell) \cdot n$, also $k \cdot \ell = 1$ und $m = n$. Damit ist die Anti-Symmetrie nachgewiesen. Um die Transitivität zu überprüfen, seien $m, n, p \in \mathbb{N}$ mit $m | n$ und $n | p$ vorgegeben. Es gibt $k, \ell \in \mathbb{N}$ mit $n = k \cdot m$ und $p = \ell \cdot n$. Wegen $p = \ell \cdot (k \cdot m) = (k \cdot \ell) \cdot m$ folgt $m | p$.

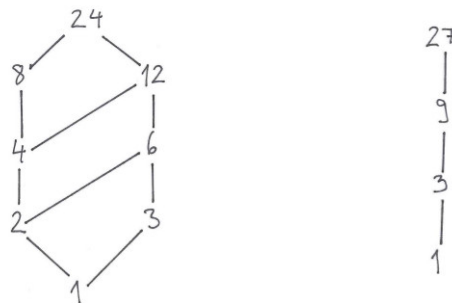
Um zu zeigen, dass $|$ keine Totalordnung ist, genügt es, zwei nicht vergleichbare Elemente $m, n \in \mathbb{N}$ anzugeben. Dies ist beispielsweise für $m = 2$ und $n = 3$ der Fall, denn weder ist 2 ein Teiler von 3, noch umgekehrt 3 ein Teiler von 2. □

- (iii) Sei X eine Menge und $\mathcal{P}(X)$ die zugehörige Potenzmenge. Dann ist die „ $\subseteq$ “-Relation auf $\mathcal{P}(X)$ eine Halbordnung. Eine Totalordnung liegt genau dann vor, wenn X aus höchstens einem Element besteht. (Man bezeichnet diese Relation auch als **Inklusionsrelation**.)

Wieder überprüfen wir zunächst die Halbordnungs-Eigenschaften. Für alle $A \in \mathcal{P}(X)$ gilt offenbar $A \subseteq A$, also ist die Relation $\subseteq$ reflexiv. Sind $A, B \in \mathcal{P}(X)$ mit $A \subseteq B$ und $B \subseteq A$ vorgegeben, dann folgt $A = B$. Also ist die Relation anti-symmetrisch. Für $A, B, C \in \mathcal{P}(X)$ folgt aus $A \subseteq B$ und $B \subseteq C$ offenbar $A \subseteq C$, also ist die Relation auch transitiv.

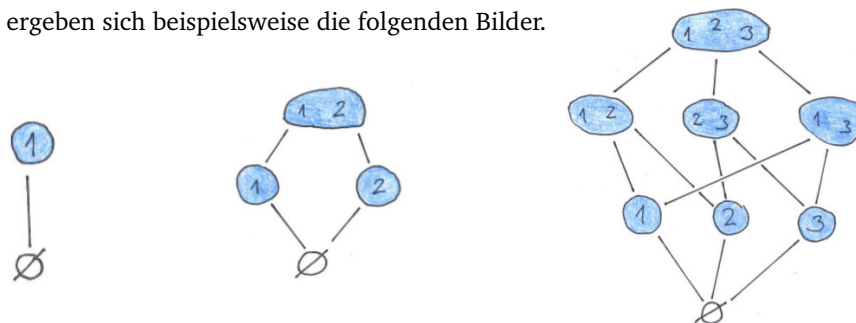
Enthält X zwei verschiedene Elemente x und y , dann ist die Relation keine Totalordnung, denn es gilt weder $\{x\} \subseteq \{y\}$ noch $\{y\} \subseteq \{x\}$. Ist $X = \emptyset$, dann gilt $\mathcal{P}(X) = \{\emptyset\}$. Es gibt also in $\mathcal{P}(X)$ gar keine zwei verschiedenen Elemente, und damit ist die Totalordnungs-Bedingung nach Definition erfüllt. Ist X einelementig, $X = \{x\}$, dann gilt $\mathcal{P}(X) = \{\emptyset, \{x\}\}$, und es gilt $\emptyset \subseteq \{x\}$. Es gibt also nur eine Möglichkeit, zwei verschiedene Elemente in $\mathcal{P}(X)$ zu wählen, und diese sind miteinander vergleichbar. Also liegt auch in diesem Fall auf $\mathcal{P}(X)$ eine Totalordnung vor. $\square$

Beschränkt man sich auf eine endliche Teilmenge von $\mathbb{N}$, zum Beispiel auf die Menge der Teiler einer festen Zahl $n \in \mathbb{N}$, dann lässt sich die Teilbarkeitsrelation graphisch darstellen. Ein von unten nach oben verlaufender Weg von einem Element x zu einem Element y soll dabei bedeuten, dass $x \mid y$ erfüllt ist. Für die Fälle $n = 24$ und $n = 27$ erhält man zum Beispiel die folgenden Bilder.



Die lineare Struktur des Graphen rechts zeigt an, dass es sich bei der Relation $\mid$ auf der Menge $\{1, 3, 9, 27\}$ um eine Totalordnung handelt. Die Teiler von 24 bilden aber nur eine Halbordnung.

Auch die Inklusionsrelation auf $\mathcal{P}(X)$ lässt sich für kleine Mengen X veranschaulichen. Für die Potenzmengen $\{1\}$, $\{1, 2\}$ und $\{1, 2, 3\}$ ergeben sich beispielsweise die folgenden Bilder.



Ist $\leq$ eine allgemeine Halbordnung auf einer beliebigen Menge X , dann ist es allgemein üblich, die folgenden abkürzenden Schreibweisen zu verwenden.

$x \geq y$	für die Aussage	$y \leq x$
$x < y$	für die Aussage	$x \leq y \wedge x \neq y$
$x > y$	für die Aussage	$y \leq x \wedge y \neq x$

Wir bemerken, dass sich die Bedingungen $x \leq y$ und $x > y$ (und ebenso die Bedingungen $x \geq y$ und $x < y$) gegenseitig ausschließen. Aus $x \leq y$ und $x > y$ würde nämlich insbesondere $x \leq y$ und $x \geq y$ und auf Grund der Anti-Symmetrie damit $x = y$ folgen, was zur Voraussetzung $x > y$ im Widerspruch steht.

(3.4) Definition Sei $(X, \leq)$ eine Halbordnung und $A \subseteq X$.

- (i) Man nennt $a \in A$ ein **größtes** (bzw. **kleinstes**) Element der Menge A , wenn $a \geq b$ (bzw. $a \leq b$) für alle $b \in A$ gilt.
- (ii) Ein Element $a \in A$ wird **maximales** (bzw. **minimales**) Element der Menge A genannt, wenn kein $b \in A$ mit $b > a$ (bzw. $b < a$) existiert.

Neben „größtes Element“ und „kleinstes Element“ sind auch die Bezeichnungen „Maximum“ und „Minimum“ gebräuchlich. Das Maximum einer Teilmenge $A \subseteq X$, sofern es existiert, wird mit $\max(A)$ bezeichnet, und ebenso das Minimum im Falle der Existenz mit $\min(A)$. Die soeben eingeführten Begriffe stehen in folgender Beziehung zueinander.

(3.5) Proposition Sei $(X, \leq)$ eine Halbordnung und $A \subseteq X$.

- (i) Es gibt höchstens ein größtes und höchstens ein kleinstes Element in A .
- (ii) Das größte (bzw. kleinste) Element von A , sofern es existiert, ist zugleich das einzige maximale (bzw. minimale) Element von A .
- (iii) Ist $(X, \leq)$ eine Totalordnung, dann sind die Begriffe „größtes Element“ und „maximales Element“ (bzw. „kleinstes Element“ und „minimales Element“) gleichbedeutend.

Beweis: zu (i) Nehmen wir an, dass a und a' beides größte Elemente von A sind. Weil a größtes Element von A und $a' \in A$ ist, gilt $a \geq a'$. Weil a' größtes Element von A und $a \in A$ ist, gilt $a' \geq a$. Aus $a \geq a'$, $a' \geq a$ und der Anti-Symmetrie der Relation $\leq$ folgt $a = a'$. Genauso zeigt man, dass es höchstens ein kleinstes Element in A geben kann.

zu (ii) Wir beschränken uns auf den Beweis der Aussage für das größte Element. Sei also a das größte Element von A . Dann muss a zugleich ein maximales Element sein, denn die Existenz eines $b \in A$ mit $b > a$ würde der Definition des größten Elements widersprechen. Nehmen wir nun an, dass b ein von a verschiedenes maximales Element der Menge A ist. Dann gilt $a \geq b$ (weil a größtes Element von A ist), zugleich aber $a \neq b$ und damit insgesamt $a > b$. Aber dies widerspricht der Maximalität des Elements b .

zu (iii) Wieder beschränken wir uns auf den Beweis der Aussage für größte und maximale Elemente. Wegen Teil (ii) genügt es zu zeigen, dass im Falle einer Totalordnung jedes maximale Element zugleich größtes Element ist. Sei also a ein maximales Element, und sei $b \in A$ beliebig. Weil $\leq$ eine Totalordnung ist, muss $a \leq b$ oder $a \geq b$ und somit auch $a < b$ oder $a \geq b$ gelten. Der Fall $a < b$ würde der Maximalität von a widersprechen. Es gilt also $a \geq b$ für alle $b \in A$, und damit ist a das größte Element von A . $\square$

Aus Teil (iii) der Proposition folgt insbesondere, dass es in Teilmengen der herkömmlichen Totalordnung $(\mathbb{R}, \leq)$ keinen Unterschied zwischen größten und maximalen bzw. kleinsten und minimalen Elementen gibt. In Halbordnungen kann es durchaus vorkommen, dass eine Teilmenge mehrere maximale oder minimale Elemente besitzt. Betrachten wir zum Beispiel in $\mathbb{N}$ mit der Teilerrelation die Teilmenge $A = \{1, 2, \dots, 10\}$, so besitzt diese sogar fünf maximale Elemente, nämlich 6, 7, 8, 9 und 10. In dieser Situation kann es dann aber wegen Teil (ii) der Proposition kein größtes Element geben.

Offenbar besitzt in einer Totalordnung $(X, \leq)$ jede zweielementige Teilmenge $\{a, b\}$ ein Maximum und ein Minimum. Denn in diesem Fall gilt $a \leq b$ oder $a \geq b$. Im ersten Fall ist a das Minimum und b das Maximum, im zweiten Fall ist es umgekehrt. Durch vollständige Induktion lässt sich leicht zeigen, dass jede nichtleere, endliche Teilmenge von A ein Minimum und ein Maximum besitzt.

(3.6) Definition Sei $(X, \leq)$ eine Halbordnung und $A \subseteq X$.

- (i) Sei Element $s \in X$ wird **obere Schranke** (bzw. **untere Schranke**) von A genannt, wenn $s \geq a$ (bzw. $s \leq a$) für alle $a \in A$ gilt.
- (ii) Wir bezeichnen mit $\mathcal{S}^+(A)$ bzw. $\mathcal{S}^-(A)$ die Menge aller oberen bzw. unteren Schranken von A .
- (iii) Das kleinste Element von $\mathcal{S}^+(A)$, sofern es existiert, wird das **Supremum** von A genannt und mit $\sup(A)$ bezeichnet. Ebenso nennt man das größte Element von $\mathcal{S}^-(A)$ im Falle der Existenz das **Infimum** von A , und bezeichnet es mit $\inf(A)$.

Ist die Menge A nichtleer, gilt aber $\mathcal{S}^+(A) = \emptyset$, dann setzt man $\sup(A) = +\infty$. Ebenso wird $\inf(A)$ im Fall $A \neq \emptyset$ und $\mathcal{S}^-(A) = \emptyset$ auf den Wert $-\infty$ gesetzt.

Auch diese Begriffe illustrieren wir anhand eines konkreten Beispiels.

(3.7) Proposition Seien $a, b \in \mathbb{R}$ mit $a < b$ und $I = \{x \in \mathbb{R} \mid a < x < b\}$. (Eine solche Teilmenge nennt man ein endliches offenes Intervall.)

- (i) Die Menge besitzt weder maximale noch minimale Elemente, also erst recht weder ein größtes noch ein kleinstes Element.
- (ii) Es gilt $\mathcal{S}^+(I) = \{x \in \mathbb{R} \mid x \geq b\}$ und $\mathcal{S}^-(I) = \{x \in \mathbb{R} \mid x \leq a\}$.
- (iii) Es gilt $\sup(I) = b$ und $\inf(I) = a$.

Beweis: Der Beweis beruht auf der folgenden allgemeinen Tatsache, deren Beweis wir nachliefern, sobald wir in der Vorlesung die angeordneten Körper definiert haben: Sind $c, d \in \mathbb{R}$ mit $c < d$, dann gelten für den Durchschnitt $e = \frac{1}{2}(c + d)$ der beiden Zahlen die Ungleichungen $c < e < d$.

zu (i) Nehmen wir an, dass $c \in I$ ein maximales Element von I ist. Dann gilt $a < c < b$. Setzen wir $c' = \frac{1}{2}(c + b)$, dann gilt $a < c < c' < b$ und somit $c' \in I$. Damit steht $c' > c$ aber im Widerspruch zur Maximalität von I . Genauso widerlegt man die Existenz minimaler Elemente. Dass es damit auch kein größtes oder kleinstes Element in I geben kann, folgt aus Teil (ii) von Proposition (3.5).

zu (ii) Wir beschränken uns auf den Beweis der Gleichung für $\mathcal{S}^+(I)$. Ist $x \geq b$, dann gilt insbesondere $x \geq b > c$ für alle $c \in I$. Somit ist x eine obere Schranke von I , also in $\mathcal{S}^+(I)$ enthalten. Setzen wir umgekehrt $x \in \mathcal{S}^+(I)$ voraus, und nehmen wir an, dass $x \geq b$ nicht erfüllt ist. Dann muss $x < b$ gelten. Setzen wir $a' = \max\{a, x\}$ und $c = \frac{1}{2}(a' + b)$, dann gilt $a \leq a' < c < b$ und somit $c \in I$. Aber damit steht $c > a' \geq x$ im Widerspruch zur Voraussetzung, dass x eine obere Schranke von I ist.

zu (iii) Nach Teil (ii) gilt $b \in \mathcal{S}^+(I)$ und $x \geq b$ für alle $x \in \mathcal{S}^+(I)$. Also ist b das kleinste Element von $\mathcal{S}^+(I)$, und es folgt $b = \sup(I)$. Der Beweis der Gleichung $a = \inf(I)$ läuft analog. $\square$

Allgemein gilt zwischen Maximum und Supremum bzw. Minimum und Infimum die folgende Beziehung.

(3.8) Satz Sei $(X, \leq)$ eine Halbordnung und $A \subseteq X$.

- (i) Das Maximum (bzw. Minimum) von A , sofern es existiert, ist zugleich das Supremum (bzw. Infimum) von A .
- (ii) Existiert das Supremum (bzw. Infimum) von A , und ist es in A enthalten, so ist es zugleich das Maximum (bzw. Minimum) von A .

Beweis: zu (i) Es genügt, die Aussage für das Maximum zu beweisen, denn für das Minimum läuft der Beweis analog. Sei also $a = \max(A)$. Aus der Definition des größten Elements folgt unmittelbar, dass $a \in A$ und $a \in \mathcal{S}^+(A)$ gilt. Ein $s \in \mathcal{S}^+(A)$ mit $s < a$ kann es nicht geben, denn wegen $a \in A$ könnte ein solches s keine obere Schranke von A sein. Also ist a zugleich das Supremum von A .

zu (ii) Auch hier beschränken wir uns wieder auf den Fall des Supremums. Gilt $s = \sup(A)$, dann gilt insbesondere $s \in \mathcal{S}^+(A)$ und somit $s \geq a$ für alle $a \in A$. Ist a zugleich in A enthalten, dann handelt es sich nach Definition um das größte Element der Menge A . $\square$

Betrachten wir zum Beispiel in der Totalordnung $(\mathbb{R}, \leq)$ eine Teilmenge der Form $I = [a, b] = \{x \in \mathbb{R} \mid a \leq x \leq b\}$ mit $a, b \in \mathbb{R}, a < b$ (ein sog. endliches abgeschlossenes Intervall), dann gilt $\max(I) = \sup(I) = b$ und $\min(I) = \inf(I) = a$.

(3.9) Definition Eine Halbordnung $(X, \leq)$, in der jede zweielementige Teilmenge $\{a, b\} \subseteq X$ ein Infimum und ein Supremum besitzt, bezeichnet man als **Verband**. Man verwendet die Bezeichnungen $a \vee b = \sup\{a, b\}$ und $a \wedge b = \inf\{a, b\}$.

Auch hier betrachten wir eine Reihe von Beispielen.

- (i) Jede Totalordnung $(X, \leq)$ ist ein Verband, mit $a \vee b = \max\{a, b\}$ und $a \wedge b = \min\{a, b\}$ für alle $a, b \in X$.
- (ii) Die natürlichen Zahlen mit der Teilerrelation bilden einen Verband. Für alle $m, n \in \mathbb{N}$ gilt jeweils $m \vee n = \text{kgV}(m, n)$ und $m \wedge n = \text{ggT}(m, n)$.
- (iii) Ist X eine beliebige Menge, dann ist $(\mathcal{P}(X), \subseteq)$ ein Verband. Für alle $A, B \in \mathcal{P}(X)$ gilt jeweils $A \vee B = A \cup B$ und $A \wedge B = A \cap B$.
- (iv) Schränkt man die Teilerrelation auf $\mathbb{N}$ auf die Teilmenge $X = \{1, 2, \dots, 10\}$ ein, so erhält man eine Halbordnung, die kein Verband mehr ist. Beispielsweise besitzt die Teilmenge $\{7, 8\}$ keine obere Schranke in X , somit erst recht keine kleinste obere Schranke, also kein Supremum.
- (v) Ebenso geht die Verbandsstruktur verloren, wenn man die Teilerrelation auf die Menge $\mathbb{N} \setminus \{56\}$ einschränkt. Es existieren dann zwar mehrere minimale obere Schranken von $\{7, 8\}$, zum Beispiel 112 oder 168, aber keine kleinste obere Schranke.

Wir kommen nun zu einer weiteren wichtigen Klasse von Relationen, den Äquivalenzrelationen.

(3.10) Definition Eine Relation $\sim$ auf einer Menge X wird **Äquivalenzrelation** genannt, wenn sie reflexiv, symmetrisch und transitiv ist.

Wir betrachten ein erstes Beispiel für eine Äquivalenzrelation. Sei X eine endliche Menge und $\mathcal{P}(X)$ die Potenzmenge von X . Dann ist durch $A \sim B \Leftrightarrow |A| = |B|$ eine Äquivalenzrelation auf $\mathcal{P}(X)$ definiert, wobei $|A|$ jeweils die Anzahl der Elemente von A bezeichnet. Diese Relation ist reflexiv, denn für alle $A \in \mathcal{P}(X)$ gilt $|A| = |A|$ und somit $A \sim A$. Sie ist symmetrisch, denn für alle $A, B \in \mathcal{P}(X)$ mit $A \sim B$ gilt $|A| = |B|$, damit auch $|B| = |A|$, und folglich $B \sim A$. Die Relation ist auch transitiv. Sind nämlich $A, B, C \in \mathcal{P}(X)$ mit $A \sim B$ und $B \sim C$ vorgegeben, dann gilt $|A| = |B|$ und $|B| = |C|$, damit auch $|A| = |C|$ und somit $A \sim C$.

Für jede natürliche Zahl n erhält man auf folgende Weise eine Relation auf der Menge $\mathbb{Z}$ der natürlichen Zahlen.

(3.11) Definition Für jedes $n \in \mathbb{N}$ sei die Relation $\equiv_n$ auf $\mathbb{Z}$ definiert durch die Festlegung

$$a \equiv_n b \Leftrightarrow n \mid (a - b) \quad \forall a, b \in \mathbb{Z}.$$

Hierbei steht $\mid$ für die Teilerrelation; die Schreibweise $n \mid (a - b)$ bedeutet also, dass n ein Teiler der ganzen Zahl $a - b$ ist. Die Relation $\equiv_n$ wird als **Kongruenzrelation modulo n** bezeichnet. Zwei ganze Zahlen $a, b \in \mathbb{Z}$ mit $a \equiv_n b$ werden auch **kongruent modulo n** genannt.

An Stelle von $a \equiv_n b$ sind auch die Schreibweisen $a \equiv b \pmod{n}$ und $a \equiv b(n)$ gebräuchlich.

(3.12) Satz Für jedes $n \in \mathbb{N}$ ist $\equiv_n$ eine Äquivalenzrelation auf $\mathbb{Z}$.

Beweis: Für alle $a \in \mathbb{Z}$ gilt offenbar $a \equiv_n a$, denn n ist stets ein Teiler von $a - a = 0$. Also ist die Relation $\equiv_n$ reflexiv. Sind $a, b \in \mathbb{Z}$ mit $a \equiv_n b$, dann gilt nach Definition $n \mid (a - b)$. Es gibt also ein $k \in \mathbb{Z}$ mit $a - b = kn$. Aber dann gilt auch $b - a = (-k)n$, also $n \mid (b - a)$, und damit auch $b \equiv_n a$. Dies zeigt, dass $\equiv_n$ eine symmetrische Relation ist.

Seien nun $a, b, c \in \mathbb{Z}$ mit $a \equiv_n b$ und $b \equiv_n c$ vorgegeben. Dann gilt $n \mid (a - b)$ und $n \mid (b - c)$, es gibt also $k, \ell \in \mathbb{Z}$ mit $a - b = kn$ und $b - c = \ell n$. Es folgt $a - c = (a - b) + (b - c) = (k + \ell)n$, also $n \mid (a - c)$ und somit $a \equiv_n c$. Dies zeigt, dass $\equiv_n$ auch transitiv ist. Insgesamt handelt es sich bei $\equiv_n$ also tatsächlich um eine Äquivalenzrelation. $\square$

Die intuitive Bedeutung der Äquivalenzrelationen auf einer Menge wird durch den folgenden Begriff besser verständlich.

(3.13) Definition Als **Zerlegung** einer Menge X bezeichnen wir eine Teilmenge $\mathcal{X} \subseteq \mathcal{P}(X)$ mit den Eigenschaften $A \neq \emptyset$ für alle $A \in \mathcal{X}$, dass für jedes $x \in X$ ein $A \in \mathcal{X}$ mit $x \in A$ existiert, und dass für alle $A, B \in \mathcal{X}$ aus $A \cap B \neq \emptyset$ jeweils $A = B$ folgt.

Die zweite Bedingung besagt also, dass X die Vereinigung aller Elemente aus $\mathcal{X}$ ist, und die dritte, dass je zwei verschiedene Mengen aus X disjunkt sind. Ist zum Beispiel $X = \{1, 2, 3, 4, 5\}$, dann ist sowohl $\{\{1, 2, 3\}, \{4, 5\}\}$ als auch $\{\{1\}, \{5\}, \{2, 3, 4\}\}$ eine Zerlegung von X . Dagegen ist $\{\{1, 2, 3\}, \{3, 4, 5\}\}$ oder $\{\{1, 3\}, \{2, 5\}\}$ oder auch $\{\emptyset, \{1, 2\}, \{3, 4\}, \{5\}\}$ keine Zerlegung von X .

(3.14) Definition Sei X eine Menge, $\sim$ eine Äquivalenzrelation auf X und $x \in X$. Dann nennt man die Teilmenge

$$[x]_{\sim} = \{y \in X \mid x \sim y\}$$

die **Äquivalenzklasse** des Elements x bezüglich $\sim$.

(3.15) Proposition Sei X eine Menge und $\sim$ eine Äquivalenzrelation auf X . Für alle $x, y \in X$ folgt aus $y \in [x]_{\sim}$ stets $[x]_{\sim} = [y]_{\sim}$. Die Äquivalenzklassen von $\sim$ bilden also eine Zerlegung der Menge X .

Beweis: Seien $x, y \in X$ mit $y \in [x]_{\sim}$ vorgegeben. Nach Definition $[x]_{\sim}$ gilt dann $x \sim y$. Zum Beweis von $[y]_{\sim} \subseteq [x]_{\sim}$ sei nun $z \in [y]_{\sim}$ vorgegeben. Dann gilt $y \sim z$. Aus $x \sim y$ und $y \sim z$ folgt $x \sim z$, auf Grund der Transitivität. Daraus wiederum folgt $z \in [x]_{\sim}$. Zum Beweis von $[x]_{\sim} \subseteq [y]_{\sim}$ sei nun $z \in [x]_{\sim}$. Dann gilt $x \sim z$. Aus $x \sim y$ und der Symmetrie von $\sim$ folgt $y \sim x$. Aus $y \sim x$ und $x \sim z$ folgt $y \sim z$. Damit ist $z \in [y]_{\sim}$ nachgewiesen. Insgesamt haben wir damit $[x]_{\sim} = [y]_{\sim}$ gezeigt.

Für jedes $x \in X$ enthält die Äquivalenzklasse $[x]_{\sim}$ auf jeden Fall das Element x , denn auf Grund der Reflexivität gilt $x \sim x$ und damit $x \in [x]_{\sim}$. Sämtliche Äquivalenzklassen sind also nicht leer, und jedes $x \in X$ ist in mindestens einer Äquivalenzklasse enthalten. Seien nun $x, y \in X$ mit $[x]_{\sim} \cap [y]_{\sim} \neq \emptyset$ vorgegeben. Dann existiert ein $z \in [x]_{\sim} \cap [y]_{\sim}$. Wie wir oben gezeigt haben, folgt daraus $[x]_{\sim} = [z]_{\sim} = [y]_{\sim}$. Damit haben wir für die Menge der Äquivalenzklassen die Eigenschaften einer Zerlegung nachgewiesen. $\square$

Umgekehrt lässt sich jeder Zerlegung eine Äquivalenzrelation zuordnen.

(3.16) Proposition Sei X eine Menge und $\mathcal{Z}$ eine Zerlegung von X . Dann ist durch die Festlegung

$$x \sim_{\mathcal{Z}} y \iff \exists A \in \mathcal{Z} : (x \in A) \wedge (y \in A) \quad \forall x, y \in X$$

eine Äquivalenzrelation auf X definiert.

Beweis: Für jedes $x \in X$ existiert nach Definition der Zerlegungen ein $A \in \mathcal{Z}$ mit $x \in A$. Die Aussage $(x \in A) \wedge (x \in A)$ ist somit erfüllt, es gilt also $x \sim_{\mathcal{Z}} x$. Dies zeigt, dass $\sim_{\mathcal{Z}}$ reflexiv ist. Seien nun $x, y \in X$ mit $x \sim_{\mathcal{Z}} y$ vorgegeben. Dann existiert ein $A \in \mathcal{Z}$ mit $(x \in A) \wedge (y \in A)$. Es gilt dann auch $(y \in A) \wedge (x \in A)$; daraus folgt $y \sim_{\mathcal{Z}} x$. Die Relation $\sim_{\mathcal{Z}}$ ist also auch symmetrisch. Seien schließlich $x, y, z \in X$ mit $x \sim_{\mathcal{Z}} y$ und $y \sim_{\mathcal{Z}} z$. Dann gibt es Menge $A, B \in \mathcal{Z}$ mit $(x \in A) \wedge (y \in A)$ und $(y \in B) \wedge (z \in B)$. Aus $y \in A \cap B$ und den Eigenschaften einer Zerlegung folgt $A = B$. Damit ist dann auch $(x \in A) \wedge (z \in A)$ erfüllt, und wir erhalten $x \sim_{\mathcal{Z}} z$. Dies zeigt, dass $\sim_{\mathcal{Z}}$ auch transitiv ist. Insgesamt ist durch $\sim_{\mathcal{Z}}$ also eine Äquivalenzrelation gegeben. $\square$

(3.17) Proposition Sei $n \in \mathbb{N}$ und $a \in \mathbb{Z}$. Dann ist die Äquivalenzklasse $[a]_n$ von a bezüglich der Relation $\equiv_n$ gegeben durch die Menge $a + n\mathbb{Z} = \{a + nk \mid k \in \mathbb{Z}\}$. Man nennt $[a]_n$ auch die **Kongruenz-** oder **Restklasse** von a modulo n .

Beweis: Zum Beweis der Inklusion $[a]_n \subseteq a + n\mathbb{Z}$ sei $c \in [a]_n$ vorgegeben. Dann gilt $a \equiv_n c$, also $n \mid (a - c)$. Es existiert also ein $k \in \mathbb{Z}$ mit $a - c = nk$. Daraus wiederum folgt $c = a + n(-k) \in a + n\mathbb{Z}$. Zum Nachweis von $a + n\mathbb{Z} \subseteq [a]_n$ sei $c \in a + n\mathbb{Z}$. Dann gilt $c = a + nk$ für ein $k \in \mathbb{Z}$. Es folgt $n(-k) = a - c$, also $n \mid (a - c)$ und somit $a \equiv_n c$. Dies wiederum ist gleichbedeutend mit $c \in [a]_n$. $\square$

Nach Proposition (3.15) bilden die Kongruenzklassen modulo einer Zahl $n \in \mathbb{N}$ also eine Zerlegung von $\mathbb{Z}$. Beispielsweise wird $\mathbb{Z}$ durch die Kongruenzklassen modulo 2 in die geraden und die ungeraden Zahlen zerlegt. Durch die Kongruenz modulo 3 erhalten wir eine Zerlegung $\mathbb{Z} = A \cup B \cup C$ von $\mathbb{Z}$ in drei Teilmengen, nämlich $A = 3\mathbb{Z} = \{\dots, -6, -3, 0, 3, 6, 9, \dots\}$, $B = 1 + 3\mathbb{Z} = \{\dots, -5, -2, 1, 4, 7, 10, \dots\}$ und $C = 2 + 3\mathbb{Z} = \{\dots, -4, -1, 2, 5, 8, 11, \dots\}$.

Wenden wir Proposition (3.15) auf die Relation $\equiv_n$ an, so erhalten wir

(3.18) Folgerung Für alle $n \in \mathbb{N}$ und $a, b \in \mathbb{Z}$ gilt die Äquivalenz

$$a \equiv_n b \iff b \in a + n\mathbb{Z} \iff a + n\mathbb{Z} = b + n\mathbb{Z}.$$

Beispielsweise gelten in $\mathbb{Z}/5\mathbb{Z}$ die Gleichungen $2 + 5\mathbb{Z} = 7 + 5\mathbb{Z} = 12 + 5\mathbb{Z} = -3 + 5\mathbb{Z}$.

Aus der Schulmathematik ist das Konzept der **Division mit Rest** bekannt: Ist $a \in \mathbb{Z}$ und $n \in \mathbb{N}$ vorgegeben, so findet man stets Elemente $q, r \in \mathbb{Z}$, so dass $a = qn + r$ und $0 \leq r < n$ erfüllt ist.

(3.19) Proposition Für jedes $n \in \mathbb{N}$ sei $\mathbb{Z}/n\mathbb{Z}$ die Menge der Kongruenzklassen modulo n .
Dann besitzt $\mathbb{Z}/n\mathbb{Z}$ genau n verschiedene Elemente, nämlich $r + n\mathbb{Z}$ mit $0 \leq r < n$.

Beweis: Zunächst zeigen wir, dass jede Kongruenzklasse mit einem dieser n Elemente übereinstimmt. Sei $a + n\mathbb{Z}$ vorgegeben, mit $a \in \mathbb{Z}$. Division mit Rest liefert $q, r \in \mathbb{Z}$ mit $a = qn + r$ und $0 \leq r < n$. Wegen $a - r = qn$ gilt $n \mid (a - r)$ und somit $a \equiv_n r$. Wir erhalten $a + n\mathbb{Z} = [a]_n = [r]_n = r + n\mathbb{Z}$, nach Proposition (3.18). Um zu zeigen, dass die angegebenen Elemente alle verschieden sind, seien $r, s \in \mathbb{Z}$ mit $0 \leq r, s < n$ vorgegeben. Nehmen wir an, es gilt $r + n\mathbb{Z} = s + n\mathbb{Z}$; zu zeigen ist, dass dann $r = s$ folgt. Aus $[r]_n = [s]_n$ folgt $s \in [r]_n$ und somit $r \equiv_n s$, also $n \mid (r - s)$. Aber wegen $0 \leq r, s < n$ ist $|r - s| < n$; somit ist $n \mid (r - s)$ nur möglich, wenn $r = s$ ist. $\square$

Statt mit $[a]_n$ oder $a + n\mathbb{Z}$ bezeichnet man die Restklasse von a auch mit $\bar{a}$, sofern n aus dem Kontext heraus bekannt ist. Nach Proposition (3.19) ist die Menge $\mathbb{Z}/7\mathbb{Z}$ beispielsweise gegeben durch

$$\mathbb{Z}/7\mathbb{Z} = \{\bar{0}, \bar{1}, \bar{2}, \bar{3}, \bar{4}, \bar{5}, \bar{6}\}.$$

§ 4. Abbildungen und Mächtigkeiten

Inhaltsübersicht

Eine Abbildung $f : X \rightarrow Y$ zwischen zwei Mengen X und Y ist gegenüber einer beliebigen Relation dadurch ausgezeichnet, dass für jedes x aus X , dem *Definitionsbereich* von f , ein eindeutig bestimmtes y aus Y , dem *Wertebereich* von f , existiert, das dann mit $f(x)$ bezeichnet wird. Man umschreibt dies mit der Formulierung, dass die Abbildung f jedem $x \in X$ ein Element $y \in Y$ „zuordnet“.

Im Allgemeinen können mehrere $x \in X$ demselben $y \in Y$ zugeordnet werden; ist dies nicht der Fall, nennt man die Abbildung *injektiv*. Ist jedes $y \in Y$ das Ziel von mindestens einer Zuordnung, wird also der gesamte Definitionsbereich durch f „abgedeckt“, spricht man von einer *surjektiven* Abbildung. Abbildungen, die injektiv und surjektiv sind, werden *bijektiv* oder auch „1-zu-1“-Abbildungen genannt. Für sie existiert eine sog. *Umkehrabbildung* in Gegenrichtung $f^{-1} : Y \rightarrow X$, die dadurch gekennzeichnet ist, dass für alle $x \in X$ und $y \in Y$ die Äquivalenz

$$y = f(x) \iff x = f^{-1}(y)$$

erfüllt ist. Mit anderen Worten, die Gleichung $y = f(x)$ kann nach x „aufgelöst“ werden.

Mit Hilfe der bijektiven Abbildungen kann auch die *Mächtigkeit* einer Menge definiert werden. Ist die Menge *endlich*, so handelt es sich dabei einfach um die Anzahl der Elemente. Mit Hilfe der vollständigen Induktion aus § 2 leiten wir eine Reihe von Rechenregeln für die Mächtigkeit endlicher Mengen her. Außerdem behandeln wir einige Grundlagen zu *unendlichen* Mengen, indem wir beispielsweise durch das Konzept der Abzählbarkeit verschiedene Stufen der Unendlichkeit unterscheiden.

Wichtige Begriffe und Sätze

- Abbildung zwischen zwei Mengen
- Definitions- und Wertebereich einer Abbildung
- Einschränkung und Komposition von Abbildungen
- Bildmengen und Urbildmengen
- injektive, surjektive und bijektive Abbildungen, Umkehrabbildung
- endliche und unendliche Mengen, Mächtigkeit einer endlichen Menge
- Rechenregeln für die Mächtigkeit endlicher Mengen
(für Vereinigung, Durchschnitt, kartesischem Produkt und Potenzmenge)
- Binomialkoeffizienten
- Gleichmächtigkeit; höchstens abzählbare, abzählbar unendliche und überabzählbare Mengen
- Familie von Elementen, Indexmenge, Folge
- Abzählbarkeitskriterien, Abzählbarkeit von $\mathbb{Q}$ als Folgerung

(4.1) Definition Seien X, Y Mengen. Eine Relation f zwischen X und Y wird **Abbildung** genannt, wenn für jedes $x \in X$ genau ein $y \in Y$ mit $(x, y) \in f$ existiert. In Formelschreibweise:

$$(i) \quad \forall x \in X : \exists y \in Y : (x, y) \in f$$

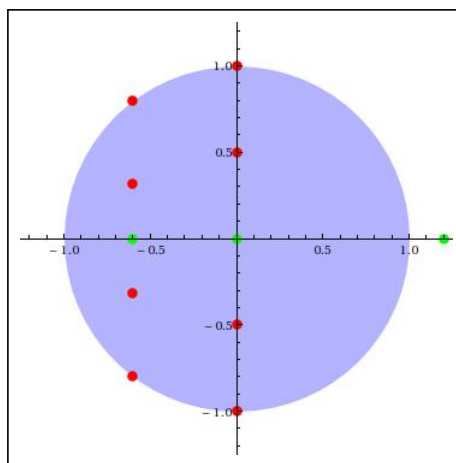
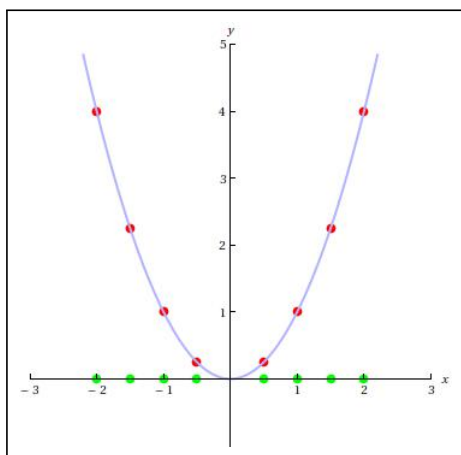
$$(ii) \quad \forall x \in X : \forall y, y' \in Y : (x, y) \in f \wedge (x, y') \in f \Rightarrow y = y'.$$

Dabei nennt man X den **Definitions-** und Y den **Wertebereich** der Abbildung.

Für gegebenes $x \in X$ bezeichnet man das eindeutig bestimmte $y \in Y$ mit der Eigenschaft $(x, y) \in R$ mit $f(x)$ und nennt es das **Bild** von x unter R .

Die Notation $f : X \rightarrow Y$ drückt aus, dass f eine Abbildung von X nach Y , also eine Abbildung mit Definitionsbereich X und Wertebereich Y ist. Die Schreibweise $x \mapsto y$ ist gleichbedeutend mit $y = f(x)$; man sagt auch, dass f dem Element x das Element y **zuordnet**.

Ob eine Relation f auf der Menge $\mathbb{R}$ der reellen Zahlen eine Abbildung ist, ob also für jedes $x \in X$ genau ein $y \in Y$ mit $(x, y) \in f$ existiert, lässt sich gut an der graphischen Darstellung von f erkennen.



Zum Beispiel ist $S = \{(x, y) \in \mathbb{R}^2 \mid y = x^2\}$ eine Abbildung, denn für jeden x -Wert (grün) gibt es genau einen zugehörigen Punkt $(x, y) \in S$ (rot), also genau ein $y \in \mathbb{R}$ mit $(x, y) \in S$. Dagegen ist $T = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 \leq 1\}$ keine Abbildung, denn für einige x -Werte (zum Beispiel für $x = 1.2$) gibt es gar kein y mit $(x, y) \in T$. Für andere x -Werte (zum Beispiel $x = 0$) gibt es dagegen gleich mehrere, sogar unendlich viele, zugehörige y -Werte, was bei einer Abbildung ebenfalls nicht zulässig ist.

Als nächstes definieren wir zwei wichtige Operationen: die Einschränkung und die Komposition von Funktionen. Die erste Operation läuft darauf hinaus, dass man den Definitionsbereich einer Funktion f verkleinert.

(4.2) Proposition Sind X, Y Mengen, $f \subseteq X \times Y$ eine Abbildung und $U \subseteq X$. Dann ist durch $f|_U = \{(x, y) \in f \mid x \in U\} = f \cap (U \times Y)$ eine Abbildung von U nach Y definiert. Wir bezeichnen sie als die **Einschränkung** von f auf die Teilmenge U .

Beweis: Um zu zeigen, dass $g \circ f = f \cap (U \times Y)$ eine Abbildung $U \rightarrow Y$ ist, sei $x \in U$ vorgegeben. Weil f eine Abbildung ist, existiert jedenfalls ein $y \in Y$ mit $(x, y) \in f$. Wegen $(x, y) \in U \times Y$ ist (x, y) auch in g enthalten. Seien nun $x \in U$ und $y, y' \in Y$ mit $(x, y) \in g$ und $(x, y') \in g$. Dann gilt auch $(x, y) \in f$ und $(x, y') \in f$. Weil f eine Abbildung ist, folgt $y = y'$. Insgesamt haben wir damit für g beide in Definition (4.1) angegebenen Bedingungen verifiziert. $\square$

(4.3) Proposition Seien X, Y, Z Mengen und $f : X \rightarrow Y, g : Y \rightarrow Z$ zwei Abbildungen. Dann ist $g \circ f = \{ (x, z) \in X \times Z \mid \exists y \in Y : (x, y) \in f \wedge (y, z) \in g \}$ eine Abbildung von X nach Z , und es gilt $(g \circ f)(x) = g(f(x))$ für alle $x \in X$. Man nennt $g \circ f$ die **Komposition** der Abbildungen f und g .

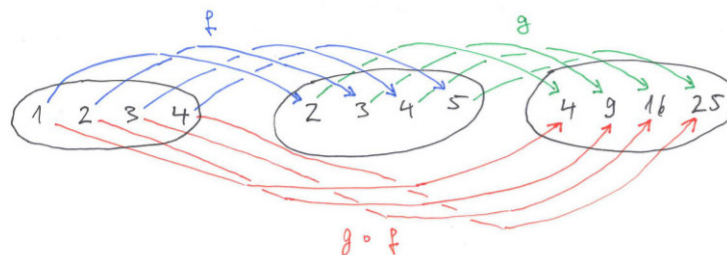
Beweis: Sei $x \in X$ vorgegeben. Setzen wir $y = f(x) \in Y$, dann gilt $(x, y) \in f$ und $(y, g(y)) \in g$ nach Definition der Bildelemente $f(x)$ und $g(y)$. Dies zeigt, dass das Paar (x, z) mit $z = g(y) = g(f(x))$ in $g \circ f$ enthalten ist.

Nehmen wir nun an, $z' \in Z$ ist ein weiteres Element mit $(x, z') \in g \circ f$ ist. Dann existiert nach Definition ein $y' \in Y$ mit $(x, y') \in f$ und $(y', z') \in g$. Weil f eine Abbildung ist, gilt $y = f(x) = y'$, und aus der Abbildungs-Eigenschaft von g und der Voraussetzung $(y, z) \in g$ und $(y', z') = (y, z') \in g$ folgt $z = z'$. Insgesamt ist damit gezeigt, dass $g \circ f$ eine Abbildung ist, und dass $(g \circ f)(x) = g(f(x))$ gilt. $\square$

Die Komposition $g \circ f$ entsteht also einfach dadurch, dass f in g „eingesetzt“ wird. Später werden wir häufig mit Abbildungen auf den reellen Zahlen arbeiten, zum Beispiel $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x + 1$ oder $g : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$. Die Komposition von f und g ergibt in diesem Fall also

$$(g \circ f)(x) = g(f(x)) = g(x + 1) = (x + 1)^2 \quad \text{für alle } x \in \mathbb{R}.$$

Man kann sich das Zusammenspiel von f, g und $g \circ f$ durch folgendes Diagramm veranschaulichen.

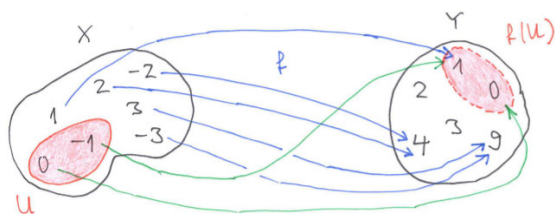


Eine Abbildung wirkt nicht nur auf einzelne Elemente, sondern auch auf Teilmengen ihres Definitions- und Wertebereichs. Die folgenden beiden Konzepte werden später bei der Formulierung der Ketten- und der Umkehrregel in der Analysis eine wichtige Rolle spielen.

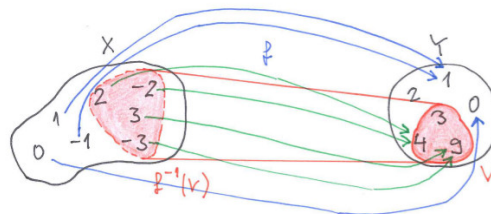
(4.4) Definition Seien $f : X \rightarrow Y$ eine Abbildung und $U \subseteq X, V \subseteq Y$.

- (i) Die Teilmenge $f(U) = \{f(x) \mid x \in U\} \subseteq Y$ wird die **Bildmenge** von U unter der Abbildung f genannt. Es handelt sich um die Elemente von Y , die man dadurch erhält, dass man f auf ein Element aus U anwendet.
- (ii) Die Teilmenge $f^{-1}(V) = \{x \in X \mid f(x) \in V\} \subseteq X$ wird die **Urbildmenge** von V unter f genannt. Sie besteht aus genau den Elementen von X , die nach V abgebildet werden.

Wir betrachten als Beispiel die Abbildung $f : X \rightarrow Y$ zwischen den Mengen $X = \{-3, -2, -1, 0, 1, 2, 3\}$ und $Y = \{0, 1, 2, 3, 4, 9\}$ gegeben durch $f(x) = x^2$ für alle $x \in X$. Seien $U \subseteq X$ und $V \subseteq Y$ gegeben durch $U = \{-1, 0\}$ und $V = \{3, 4, 9\}$.



Bestimmung der Bildmenge $f(U)$



Bestimmung der Urbildmenge $f^{-1}(V)$

Nach Definition besteht $f(U)$ aus allen Elementen, die man durch Quadrierung von Elementen aus $U = \{-1, 0\}$ erhält, also $(-1)^2 = 1$ und $0^2 = 0$. Die Urbildmenge $f^{-1}(V)$ enthält alle Elemente $x \in X$, deren Quadrat in $V = \{3, 4, 9\}$ liegt. Wegen $(-2)^2 = 2^2 = 4$ und $(-3)^2 = 3^2 = 9$ sind dies die Zahlen $-3, -2, 2, 3$.

(4.5) Proposition Sei $f : X \rightarrow Y$ eine Abbildung, $U \subseteq X$ und $V \subseteq Y$. Dann gilt

$$(i) f(f^{-1}(V)) \subseteq V \quad (ii) U \subseteq f^{-1}(f(U))$$

Beweis: zu (i) Ist $y \in f(f^{-1}(V))$, dann gibt es nach Definition der Bildmenge ein $x \in f^{-1}(V)$ mit $y = f(x)$. Aus $x \in f^{-1}(V)$ folgt $f(x) \in V$, insgesamt also $y = f(x) \in V$.

zu (ii) Sei $x \in U$ vorgegeben. Dann gilt $f(x) \in f(U)$. Das Element x wird also auf ein Element aus $f(U)$ abgebildet. Nach Definition der Urbildmenge folgt daraus unmittelbar $x \in f^{-1}(f(U))$. $\square$

Anhand passender Beispiele werden wir uns in den Übungen klarmachen, dass die beiden Inklusionen in der Proposition im Allgemeinen keine Gleichungen sind, es braucht also weder $f(f^{-1}(V)) = V$ noch $f^{-1}(f(U)) = U$ zu gelten. Die Operationen „Bildmenge“ und „Urbildmenge“ heben sich also nicht immer gegenseitig auf.

(4.6) Definition Sei $f : X \rightarrow Y$ eine Abbildung.

- (i) Wenn für alle x_1, x_2 aus $f(x_1) = f(x_2)$ jeweils $x_1 = x_2$ folgt, dann nennt man die Abbildung **injektiv**.
- (ii) Wenn es für jedes $y \in Y$ ein $x \in X$ mit $f(x) = y$ gibt, dann wird f **surjektiv** genannt.
- (iii) Eine Abbildung f , die sowohl injektiv als auch surjektiv ist, bezeichnet man als **bijektiv**.

Die drei soeben definierten Eigenschaften von Abbildungen lassen sich auch mit Hilfe der Urbildmengen charakterisieren. Eine Abbildung ist genau dann injektiv, wenn $f^{-1}(\{y\})$ für jedes $y \in Y$ **höchstens** ein Element enthält. Zum

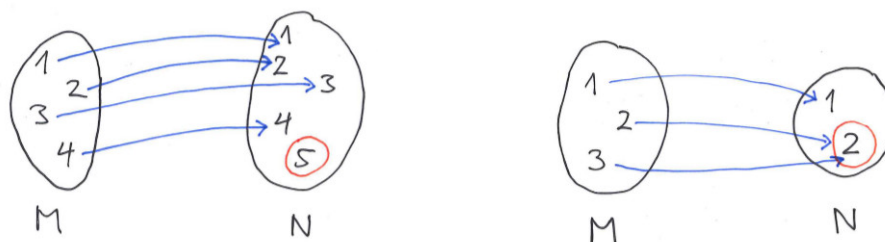
Beweis setzen wir die Injektivität voraus und geben uns ein beliebiges $y \in Y$ vor. Sind nun $x_1, x_2 \in f^{-1}(\{y\})$, dann gilt $f(x_1) = y = f(x_2)$, und aus der Injektivität folgt $x_1 = x_2$. Dies zeigt, dass $f^{-1}(\{y\})$ keine zwei verschiedenen Elemente enthält. Setzen wir umgekehrt voraus, dass $f^{-1}(\{y\})$ für jedes $y \in Y$ höchstens ein Element enthält, und seien $x_1, x_2 \in X$ mit $f(x_1) = f(x_2)$. Setzen wir $y = f(x_1)$, dann sind x_1, x_2 beides Elemente von $f^{-1}(\{y\})$. Weil aber $f^{-1}(\{y\})$ nach Voraussetzung höchstens einelementig ist, muss $x_1 = x_2$ gelten. Damit ist die Injektivität bewiesen.

Auf ähnliche Weise zeigt man, dass die Surjektivität gleichbedeutend damit ist, dass $f^{-1}(\{y\})$ für jedes $y \in Y$ aus **mindestens** einem Element besteht. Ebenfalls zur Surjektivität äquivalent ist die Gleichung $f(X) = Y$, denn nach Definition besteht $f(X)$ genau aus den Elementen $y \in Y$, für die ein $x \in X$ mit $f(x) = y$ existiert.

Aus den Feststellungen zur Injektivität und Surjektivität ergibt sich unmittelbar, dass eine Abbildung $f : X \rightarrow Y$ genau dann bijektiv ist, wenn die Menge $f^{-1}(\{y\})$ für jedes $y \in Y$ jeweils **genau** ein Element enthält.

Wieder schauen wir uns die neuen Begriffe anhand einer Reihe von Beispielen an.

- (i) Sei $M = \{1, 2, 3, 4\}$ und $N = \{1, 2, 3, 4, 5\}$. Die Abbildung $f : M \rightarrow N$ mit $f(a) = a$ für $1 \leq a \leq 4$ ist injektiv. Dafür müssen wir zeigen, dass die Aussage $\forall x_1, x_2 \in M : f(x_1) = f(x_2) \Rightarrow x_1 = x_2$ gültig ist. Seien also $a_1, a_2 \in M$ mit $f(a_1) = f(a_2)$ vorgegeben. Dann gilt $a_1 = f(a_1) = f(a_2) = a_2$, also ist die Implikation $f(a_1) = f(a_2) \Rightarrow a_1 = a_2$ tatsächlich für alle $a_1, a_2 \in M$ erfüllt. Die Abbildung ist aber nicht surjektiv, denn es gibt kein $a \in M$ mit $f(a) = 5$.



- (ii) Sei $M = \{1, 2, 3\}$ und $N = \{1, 2\}$. Dann ist die Abbildung $f : M \rightarrow N$ gegeben durch $1 \mapsto 1, 2 \mapsto 2, 3 \mapsto 2$ zwar surjektiv, aber nicht injektiv. Zwar gibt es für jedes $b \in N = \{1, 2\}$ ein $a \in M$ mit $f(a) = b$ ($f(1) = 1, f(2) = 2$), aber die Implikation $f(a_1) = f(a_2) \Rightarrow a_1 = a_2$ ist beispielsweise für $a_1 = 2, a_2 = 3$ nicht erfüllt.
- (iii) Die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^2$ ist weder injektiv noch surjektiv. Die Implikation $f(a_1) = f(a_2) \Rightarrow a_1 = a_2$ ist beispielsweise für $a_1 = -1, a_2 = 1$ nicht erfüllt, und es gibt kein $a \in \mathbb{R}$ mit $f(a) = -1$.
- (iv) Die Abbildung $f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x + 1$ ist bijektiv. Zunächst zeigen wir die Injektivität. Sind $a_1, a_2 \in \mathbb{R}$ beliebig vorgegeben, dann gelten die Implikationen

$$f(a_1) = f(a_2) \Rightarrow a_1 + 1 = a_2 + 1 \Rightarrow a_1 = a_2,$$

also ist $\forall x_1, x_2 \in \mathbb{R} : f(x_1) = f(x_2) \Rightarrow x_1 = x_2$ wahr. Ist $b \in \mathbb{R}$ beliebig vorgegeben, dann setzen wir $a = b - 1$ und erhalten $f(a) = f(b - 1) = (b - 1) + 1 = b$. Für jedes $b \in \mathbb{R}$ gibt es also ein $a \in \mathbb{R}$ mit $f(a) = b$, d.h. die Aussage $\forall y \in \mathbb{R} : \exists x \in \mathbb{R} : f(x) = y$ ist erfüllt.

- (v) Für jede Menge X nennt man $\text{id}_X : X \rightarrow X, x \mapsto x$ die **identische Abbildung** oder **Identität** auf der Menge X . Sie ist offenbar ebenfalls bijektiv.

(4.7) Satz Die Komposition zweier injektiver (bzw. surjektiver, bijektiver) Abbildungen ist injektiv (bzw. surjektiv, bijektiv).

Beweis: Seien X, Y, Z Mengen und $f : X \rightarrow Y, g : Y \rightarrow Z$ Abbildungen. Zunächst setzen wir voraus, dass f und g injektiv sind und beweisen die Injektivität von $g \circ f$. Seien dazu $x_1, x_2 \in X$ mit $(g \circ f)(x_1) = (g \circ f)(x_2)$ vorgegeben. Nach Definition der Komposition $\circ$ ist dies gleichbedeutend mit $g(f(x_1)) = g(f(x_2))$. Weil g injektiv ist, folgt daraus $f(x_1) = f(x_2)$. Weil auch f injektiv ist, erhalten wir $x_1 = x_2$. Damit ist die Injektivität von $g \circ f$ bewiesen.

Nun setzen wir voraus, dass f und g surjektiv sind, und beweisen die Surjektivität von $g \circ f$. Sei $z \in Z$ vorgegeben. Zu zeigen ist, dass ein $x \in X$ mit $(g \circ f)(x) = z$ existiert. Weil g surjektiv ist, gibt es ein $y \in Y$ mit $g(y) = z$. Weil auch f surjektiv ist, existiert ein $x \in X$ mit $f(x) = y$. Insgesamt gilt also $(g \circ f)(x) = g(f(x)) = g(y) = z$. Damit ist die Surjektivität von $g \circ f$ bewiesen.

Setzen wir nun voraus, dass f und g bijektiv sind. Dann sind f und g insbesondere injektiv, und wie wir im ersten Absatz gezeigt haben, folgt daraus die Injektivität von $g \circ f$. Die Abbildung f und g sind auch beide surjektiv. Wie im zweiten Absatz gezeigt, folgt daraus die Surjektivität von $g \circ f$. Als injektive und surjektive Abbildung ist $g \circ f$ also insgesamt bijektiv. $\square$

Oft werden wir auch die folgende Charakterisierung injektiver, surjektiver und bijektiver Abbildungen verwenden.

(4.8) Satz Seien X, Y nichtleere Mengen und $f : X \rightarrow Y$ eine Abbildung.

- (i) Es ist f genau dann injektiv, wenn eine Abbildung $g : Y \rightarrow X$ mit $g \circ f = \text{id}_X$ existiert.
- (ii) Sie ist genau dann surjektiv, wenn es ein $g : Y \rightarrow X$ mit $f \circ g = \text{id}_Y$ gibt.
- (iii) Sie ist bijektiv genau dann, wenn ein $g : Y \rightarrow X$ mit den Eigenschaften $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$ existiert. Die Abbildung g mit diesen beiden Eigenschaften ist dann eindeutig bestimmt. Man nennt sie die **Umkehrabbildung** von f und bezeichnet sie mit f^{-1} .

Beweis: zu (i) „ $\Rightarrow$ “ Sei $f : X \rightarrow Y$ eine injektive Abbildung und $x_0 \in X$ ein beliebig gewähltes Element. Gibt es für $y \in Y$ ein Urbild von $x \in X$, so ist dieses eindeutig bestimmt, und wir definieren $g(y) = x$. Besitzt y dagegen kein Urbild, dann setzen wir $g(y) = x_0$. Ist nun $x \in X$ und $y = f(x)$, dann ist x das eindeutig bestimmte Urbild von y , und nach Definition von g gilt $(g \circ f)(x) = g(f(x)) = g(y) = x = \text{id}_X(x)$. Also besitzt g die gewünschte Eigenschaft $g \circ f = \text{id}_X$.

„ $\Leftarrow$ “ Sei $f : X \rightarrow Y$ eine Abbildung und $g : Y \rightarrow X$ mit $g \circ f = \text{id}_X$. Wir müssen zeigen, dass f injektiv ist. Seien dazu $x_1, x_2 \in X$ Elemente mit $f(x_1) = f(x_2)$. Dann gilt

$$x_1 = \text{id}_X(x_1) = (g \circ f)(x_1) = g(f(x_1)) = g(f(x_2)) = (g \circ f)(x_2) = \text{id}_X(x_2) = x_2.$$

Also ist f tatsächlich injektiv.

zu (ii) „ $\Rightarrow$ “ Sei $f : X \rightarrow Y$ eine surjektive Abbildung. Für jedes $y \in Y$ wählen wir ein beliebiges Urbild $x_y \in f^{-1}(\{y\})$ und definieren $g(y) = x_y$. Für jedes $y \in Y$ gilt dann $(f \circ g)(y) = f(g(y)) = f(x_y) = y = \text{id}_Y(y)$, also besitzt g die gewünschte Eigenschaft.

„ $\Leftarrow$ “ Sei $f : X \rightarrow Y$ eine Abbildung und $g : Y \rightarrow X$ mit $f \circ g = \text{id}_Y$. Um die Surjektivität von f nachzuweisen, müssen wir zeigen, dass es für jedes $y \in Y$ ein $x \in X$ mit $f(x) = y$ existiert. Ein solches Element x ist durch $x = g(y)$ gegeben, denn es gilt $f(g(y)) = (f \circ g)(y) = \text{id}_Y(y) = y$.

zu (iii) „ $\Leftarrow$ “ Sei $f : X \rightarrow Y$ eine Abbildung und $g : Y \rightarrow X$ mit $g \circ f = \text{id}_X$ und $f \circ g = \text{id}_Y$. Dann ist f nach (i) injektiv, nach (ii) surjektiv, insgesamt also bijektiv.

„ $\Rightarrow$ “ Sei $f : X \rightarrow Y$ eine bijektive Abbildung. Dann gibt es nach (i) eine Abbildung $g_1 : Y \rightarrow X$ mit $g_1 \circ f = \text{id}_X$ und nach (ii) eine Abbildung $g_2 : Y \rightarrow X$ mit $f \circ g_2 = \text{id}_Y$. Wir zeigen, dass $g_1 = g_2$ gilt. Für gegebenes $y \in Y$ erhalten wir auf Grund unserer Voraussetzungen

$$g_1(y) = g_1(\text{id}_Y(y)) = g_1((f \circ g_2)(y)) = g_1(f(g_2(y))) = (g_1 \circ f)(g_2(y)) = \text{id}_X(g_2(y)) = g_2(y).$$

Also gilt tatsächlich $g_1 = g_2$, d.h. $g = g_1$ ist eine Abbildung mit den beiden gewünschten Eigenschaften. Sei nun $h : Y \rightarrow X$ eine weitere Abbildung mit $h \circ f = \text{id}_X$ und $f \circ h = \text{id}_Y$. Aus $g \circ f = \text{id}_X$ und $f \circ h = \text{id}_Y$ folgt dann, wie soeben gezeigt, die Identität $g = h$. Also ist g eindeutig bestimmt. $\square$

Im Folgenden bezeichnet $M_n = \{1, 2, 3, \dots, n\}$ für jedes $n \in \mathbb{N}$ die Menge der natürlichen Zahlen von 1 bis n . Außerdem setzen wir $M_0 = \emptyset$.

(4.9) Definition Sei $n \in \mathbb{N}_0$. Man sagt, eine Menge A besteht aus n Elementen oder hat die **Mächtigkeit** n , falls eine bijektive Abbildung $\varphi : M_n \rightarrow A$ existiert. Wir schreiben dann $|A| = n$.

Darauf aufbauend können wir definieren

(4.10) Definition Eine Menge A ist **endlich**, falls ein $n \in \mathbb{N}_0$ mit $|A| = n$ existiert. Ansonsten bezeichnen wir die Menge A als **unendlich**.

Wir müssen sicherstellen, dass unsere Definition der Mächtigkeit einer endlichen Menge eindeutig ist, dass also nicht $|A| = m$ und $|A| = n$ für zwei verschiedene Zahlen $m, n \in \mathbb{N}_0$ gilt. Dies erfordert ein wenig Aufwand.

(4.11) Lemma Sei A eine beliebige Menge, und seien $a, b \in A$ zwei verschiedene Elemente. Dann ist die Abbildung $\tau_{ab} : A \rightarrow A$ gegeben durch

$$\tau_{ab}(x) = \begin{cases} b & \text{falls } x = a \\ a & \text{falls } x = b \\ x & \text{sonst} \end{cases} \quad \text{eine Bijektion.}$$

(Die Abbildung τ_{ab} vertauscht die beiden Elemente a und b miteinander, alle übrigen Elemente werden auf sich selbst abgebildet.)

Beweis: Zunächst beweisen wir die Surjektivität. Sei $y \in A$ vorgegeben. Ist $y = a$, dann gilt $\tau_{ab}(b) = y$. Im Fall $y = b$ ist $\tau_{ab}(a) = y$, und im verbleibenden Fall $y \notin \{a, b\}$ gilt $\tau_{ab}(y) = y$. Also liegt y auf jeden Fall in der Bildmenge der Abbildung τ_{ab} . Zum Nachweis der Injektivität seien $u, v \in A$ mit $\tau_{ab}(u) = \tau_{ab}(v)$ vorgegeben. Ist $\tau_{ab}(u) = \tau_{ab}(v) = a$, dann muss $u = v = b$ gelten. Im Fall $\tau_{ab}(u) = \tau_{ab}(v) = b$ gilt $u = v = a$. Ist schließlich $\tau_{ab}(u)$ nicht in $\{a, b\}$ enthalten, dann folgt $u = \tau_{ab}(u) = \tau_{ab}(v) = v$. $\square$

Für den Beweis der nächsten Aussage bemerken wir vorweg, dass die vollständige Induktion aus § 2 statt über $\mathbb{N}$ auch über $\mathbb{N}_0 = \mathbb{N} \cup \{0\}$ geführt werden kann, wenn man den Induktionsanfang bei $n = 0$ ansetzt.

(4.12) Proposition Sei $n \in \mathbb{N}_0$. Dann ist jede injektive Abbildung $M_n \rightarrow M_n$ auch surjektiv.

Beweis: Wir beweisen die Aussage durch vollständige Induktion über $n \in \mathbb{N}_0$. Die einzige Abbildung von M_0 nach M_0 ist wegen $M_0 = \emptyset$ die leere Menge, und diese ist nach Definition sowohl injektiv als auch surjektiv (also bijektiv). Damit ist die Aussage für $n = 0$ bewiesen.

Sei nun $n \in \mathbb{N}_0$ vorgegeben und $\psi : M_{n+1} \rightarrow M_{n+1}$ eine injektive Abbildung. Zu zeigen ist, dass ψ auch surjektiv ist. Zunächst betrachten wir den Fall, dass $\psi(n+1) = n+1$ ist. Dann gilt $\psi(M_n) \subseteq M_n$. Denn wäre dies nicht der Fall, dann gäbe es ein $k \in M_n$ mit $\psi(k) = n+1$, was aber wegen $\psi(k) = n+1 = \psi(n+1)$ im Widerspruch zur Injektivität von ψ stehen würde. Mit ψ ist auch die Einschränkung $\psi|_{M_n}$ injektiv. Es handelt sich also um eine injektive Abbildung $M_n \rightarrow M_n$; laut Induktionsvoraussetzung ist diese auch surjektiv. Daraus folgt $\psi(M_n) = M_n$. Wegen $\psi(n+1) = n+1$ ist damit insgesamt gezeigt, dass jedes Element in M_{n+1} von ψ getroffen wird, die Abbildung ψ also surjektiv ist. Damit ist die Betrachtung dieses Falls abgeschlossen.

Betrachten wir nun den Fall, dass $\psi(n+1) = k$ mit $k \in M_n$ gilt. Weil $\tau_{k,n+1}$ nach Lemma (4.11) bijektiv ist, ist nach Satz (4.7) mit ψ auch die Abbildung $\tilde{\psi} = \tau_{k,n+1} \circ \psi$ injektiv; darüber hinaus gilt

$$\tilde{\psi}(n+1) = (\tau_{k,n+1} \circ \psi)(n+1) = \tau_{k,n+1}(\psi(n+1)) = \tau_{k,n+1}(k) = n+1.$$

Wie im vorherigen Absatz gezeigt, folgt daraus, dass $\tilde{\psi}$ surjektiv ist. Aber damit ist auch $\psi = \tau_{k,n+1}^{-1} \circ \tilde{\psi}$ surjektiv. Damit ist der Induktionsschritt abgeschlossen. $\square$

Aus Proposition (4.12) folgt nun in der Tat die Eindeutigkeit von $|A|$ für eine endliche Menge A . Denn nehmen wir an, es gäbe $m, n \in \mathbb{N}_0$ mit $m < n$ und der Eigenschaft, dass sowohl $|A| = m$ als auch $|A| = n$ erfüllt ist. Dann gäbe es Bijektionen $\varphi : M_m \rightarrow A$ und $\psi : M_n \rightarrow A$. Nach Satz (4.7) wäre dann durch $\alpha = \varphi^{-1} \circ \psi$ eine bijektive Abbildung $M_n \rightarrow M_m$ gegeben. Wegen $M_m \subseteq M_n$ können wir α als injektive Abbildung $M_n \rightarrow M_n$ auffassen. Wegen $\alpha(M_n) = M_m \subsetneq M_n$ ist α als Abbildung $M_n \rightarrow M_n$ jedoch nicht surjektiv. Wir haben also eine injektive, nicht surjektive Abbildung $M_n \rightarrow M_n$ konstruiert. Aber die Existenz einer solchen Abbildung ist nach Proposition (4.12) ausgeschlossen. Also war unsere Annahme falsch, es kann nicht gleichzeitig $|A| = m$ und $|A| = n$ gelten.

(4.13) Proposition Zwei endliche Mengen A, B haben genau dann dieselbe Mächtigkeit, wenn eine Bijektion $A \rightarrow B$ existiert.

Beweis: „ $\Rightarrow$ “ Sei $n \in \mathbb{N}_0$ mit $|A| = n = |B|$. Die Gleichung $|A| = n$ bedeutet, dass eine bijektive Abbildung $\varphi : M_n \rightarrow A$ existiert. Aus $|B| = n$ folgt, dass es eine Bijektion $\psi : M_n \rightarrow B$ gibt. Nach Satz (4.7) ist durch $\psi \circ \varphi^{-1}$ eine Bijektion von A nach B gegeben.

„ $\Leftarrow$ “ Weil A endlich ist, gibt es ein $n \in \mathbb{N}_0$ und eine Bijektion $\varphi : M_n \rightarrow A$. Außerdem existiert nach Voraussetzung eine Bijektion $\psi : A \rightarrow B$. Somit ist $\psi \circ \varphi$ eine Bijektion $M_n \rightarrow B$, und daraus folgt $|B| = n = |A|$. $\square$

(4.14) Proposition Eine Menge A ist genau dann unendlich, wenn eine injektive Abbildung $\mathbb{N} \rightarrow A$ existiert.

Beweis: „ $\Leftarrow$ “ Nehmen wir an, es gibt eine injektive Abbildung $\psi : \mathbb{N} \rightarrow A$, obwohl A endlich ist. Setzen wir $|A| = n$, dann existiert also eine bijektive Abbildung $\varphi : M_n \rightarrow A$. Durch Einschränkung von ψ auf M_{n+1} erhalten wir eine injektive Abbildung $M_{n+1} \rightarrow A$. Durch $\alpha = \varphi^{-1} \circ (\psi|_{M_{n+1}})$ ist dann eine injektive Abbildung $M_{n+1} \rightarrow M_n$ gegeben. Aufgefasst als Abbildung $M_{n+1} \rightarrow M_{n+1}$ ist diese injektiv, aber nicht surjektiv wegen $\alpha(M_{n+1}) = M_n \subsetneq M_{n+1}$. Da nach Proposition (4.12) eine solche Abbildung nicht existiert, war unsere Annahme falsch. Wenn eine injektive Abbildung $\psi : \mathbb{N} \rightarrow A$ existiert, muss A also unendlich sein.

„ $\Rightarrow$ “ Nun setzen wir voraus, dass A unendlich ist. Wir konstruieren eine Abbildung $\psi : \mathbb{N} \rightarrow A$, indem wir die Bilder $\psi(n)$ der Reihe nach definieren. Zunächst wählen wir ein beliebiges Element $a \in A$ und setzen $\psi(1) = a$. Diese Abbildung ist offenbar injektiv. Sei nun $n \in \mathbb{N}$, und nehmen wir an, dass ψ auf der Teilmenge $M_n \subseteq \mathbb{N}$ bereits definiert und dort injektiv ist. Wäre $\psi(M_n) = A$, dann hätten wir eine Bijektion zwischen M_n und A . Die Menge A wäre dann endlich, im Widerspruch zur Annahme.

So aber können wir ein neues Element $a \in A \setminus \psi(M_n)$ wählen und $\psi(n+1) = a$ setzen. Dann ist ψ auf M_{n+1} weiterhin injektiv, denn wegen der Injektivität von $\psi|_{M_n}$ ist $\psi(k) = \psi(\ell)$ für $k < \ell$ und $k, \ell \in M_n$ ausgeschlossen. Wegen $\psi(n+1) = a \notin \psi(M_n)$ ist $k \in M_n$ und $\ell = n+1$ ebenfalls unmöglich. Wir erhalten so eine Abbildung ψ , die auf ganz $\mathbb{N}$ definiert ist. Auch diese ist injektiv. Wäre nämlich $\psi(k) = \psi(\ell)$ für zwei $k, \ell \in \mathbb{N}$ mit $k < \ell$, dann würde sich ein Widerspruch zur Injektivität von $\psi|_{M_\ell}$ ergeben. $\square$

(4.15) Satz (Rechenregeln für Mächtigkeiten)

- (i) Sind A und B endliche **disjunkte** Mengen, ist also $A \cap B = \emptyset$, dann gilt $|A \cup B| = |A| + |B|$.
- (ii) Ist B endlich und $A \subseteq B$, dann gilt $|A| \leq |B|$ und $|B \setminus A| = |B| - |A|$.
- (iii) Sind A und B beliebige endliche Mengen, dann gilt $|A \cup B| = |A| + |B| - |A \cap B|$ und $|A \times B| = |A| \cdot |B|$.
- (iv) Für jede endliche Menge A gilt $|\mathcal{P}(A)| = 2^{|A|}$.

Ist A eine endliche Menge, dann ist also $\mathcal{P}(A)$ und jede Teilmenge von A endlich.

Beweis: zu (i) Sei $m = |A|$ und $n = |B|$. Dann gibt es Bijektionen $\varphi : M_m \rightarrow A$ und $\psi : M_n \rightarrow B$. Wir definieren nun eine Abbildung $\alpha : M_{m+n} \rightarrow A \cup B$ durch $\alpha(k) = \varphi(k)$ für $1 \leq k \leq m$ und $\alpha(k) = \psi(k - m)$ für $m + 1 \leq k \leq n$. Wenn wir zeigen können, dass α bijektiv ist, dann folgt daraus $|A \cup B| = m + n = |A| + |B|$.

Zum Nachweis der Injektivität seien $k, \ell \in M_{m+n}$ mit $\alpha(k) = \alpha(\ell)$ vorgegeben. Ist $k \in M_m$, dann muss auch $\ell \in M_m$ gelten, denn ansonsten wäre $\alpha(k) = \alpha(\ell)$ ein Element von $A \cap B$, was wegen $A \cap B = \emptyset$ ausgeschlossen ist. Nach Definition der Abbildung α folgt daraus $\varphi(k) = \alpha(k) = \alpha(\ell) = \varphi(\ell)$, und weil φ injektiv ist, erhalten wir $k = \ell$. Ist $k > m$, dann folgt wegen $A \cap B = \emptyset$ ebenso $\ell > m$. Wir erhalten $\alpha(k) = \psi(k - m) = \psi(\ell - m) = \alpha(\ell)$ und wiederum $k = \ell$, diesmal auf Grund der Injektivität von ψ .

Zum Nachweis der Surjektivität sei $x \in A \cup B$ vorgegeben. Dann gilt $x \in A$ oder $x \in B$. Ist $x \in A$, dann gibt es ein $k \in M_m$ mit $\varphi(k) = x$. Daraus folgt $\alpha(k) = x$. Gilt dagegen $x \in B$, so existiert ein $k \in M_n$ mit $\psi(k) = x$, und wir erhalten $\alpha(k + m) = x$. Damit ist die Surjektivität bewiesen, insgesamt ist α also bijektiv.

zu (ii) Sei $n = |B|$. Zum Beweis von $|A| \leq n$ nehmen wir an, dass A unendlich ist oder zumindest $|A| \geq n + 1$ gilt. Im ersten Fall gibt es nach Proposition (4.14) eine injektive Abbildung $\mathbb{N} \rightarrow A$, im zweiten eine bijektive Abbildung $M_r \rightarrow A$ für ein $r \geq n + 1$. In beiden Fällen können wir die Abbildung zu einer injektiven Abbildung $M_{n+1} \rightarrow A$ einschränken, die wir wegen $A \subseteq B$ auch als injektive Abbildung $\varphi : M_{n+1} \rightarrow B$ betrachten können. Wegen $|B| = n$ gibt es nun eine Bijektion $\psi : M_n \rightarrow B$. Durch $\psi^{-1} \circ \varphi$ ist dann eine injektive Abbildung $M_{n+1} \rightarrow M_n$ definiert. Fassen wir diese als Abbildung $\alpha : M_{n+1} \rightarrow M_{n+1}$ auf, so ist α zwar injektiv, wegen $\alpha(M_{n+1}) = M_n \subsetneq M_{n+1}$ aber nicht surjektiv. Die Existenz einer solchen Abbildung ist durch Proposition (4.12) ausgeschlossen. Also war unsere Annahme falsch, und $|A| \leq n$ ist bewiesen. Weil die Menge B disjunkt in A und $B \setminus A$ zerlegt werden kann, gilt nach Teil (i) $|B| = |A| + |B \setminus A|$, also $|B \setminus A| = |B| - |A|$.

zu (iii) Zum Beweis der ersten Gleichung zerlegen wir $|A \cup B|$ disjunkt in die Teilmengen $A \cap B$, $A \setminus (A \cap B)$ und $B \setminus (A \cap B)$. Durch Anwendung von (i) und (ii) erhalten wir

$$\begin{aligned} |A \cup B| &= |A \cap B| + |A \setminus (A \cap B)| + |B \setminus (A \cap B)| = |A \cap B| + (|A| - |A \cap B|) + (|B| - |A \cap B|) \\ &= |A| + |B| - |A \cap B|. \end{aligned}$$

Die Gleichung $|A \times B| = |A| \cdot |B|$ beweisen wir durch vollständige Induktion über $n = |B|$. Ist $n = 0$, dann gilt $B = \emptyset$ und $A \times B = \emptyset$, also $|A \times B| = 0 = |A| \cdot 0 = |A| \cdot |B|$. Ist $n = 1$, dann gilt $B = \{b\}$ für ein $b \in B$. Wir bemerken, dass durch $A \rightarrow A \times B$, $a \mapsto (a, b)$ eine Bijektion gegeben ist. Denn aus $(a_1, b) = (a_2, b)$ folgt $a_1 = a_2$, also ist die Abbildung injektiv. Andererseits hat jedes Element in $A \times B$ die Form (a, b) für ein $a \in A$, stimmt also mit dem Bild von a überein. Daraus folgt die Surjektivität. Insgesamt ist die Abbildung bijektiv. Mit Proposition (4.13) erhalten wir $|A \times B| = |A| = |A| \cdot |B|$.

Sei nun $n \in \mathbb{N}$ vorgegeben, und setzen wir die Aussage für dieses n voraus. Sei $|B| = n+1$, $b \in B$ ein beliebig gewähltes Element und $B' = B \setminus \{b\}$. Nach (i) gilt $|B'| = |B| - 1 = n$, und die Induktionsvoraussetzung liefert $|A \times B'| = |A| \cdot n$. Weil $A \times B$ sich disjunkt in die Teilmengen $A \times B'$ und $A \times \{b\}$ zerlegen lässt, gilt

$$|A \times B| = |A \times B'| + |A \times \{b\}| = |A| \cdot n + |A| = |A|(n+1) = |A| \cdot |B|.$$

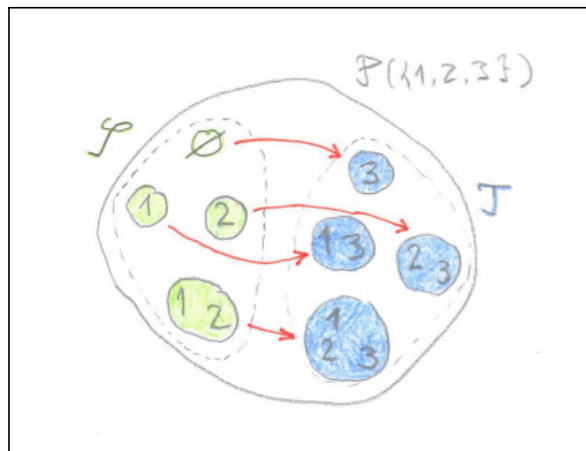
zu (iv) Der Beweis erfolgt durch vollständige Induktion über $n = |A|$. Ist $n = 0$, dann gilt $A = \emptyset$. Die leere Menge besitzt nur eine Teilmenge, nämlich $\emptyset$. Es gilt also $\mathcal{P}(A) = \{\emptyset\}$ und somit $|\mathcal{P}(A)| = 1 = 2^0$. Sei nun $n \in \mathbb{N}_0$, und setzen wir die Aussage für n voraus. Sei nun A eine $(n+1)$ -elementige Menge. Zu zeigen ist $|\mathcal{P}(A)| = 2^{n+1}$. Dazu wählen wir ein beliebiges Element $a \in A$ und setzen $A' = A \setminus \{a\}$. Dann gilt $|A'| = n$, und nach Induktionsvoraussetzung gilt $|\mathcal{P}(A')| = 2^n$. Wir betrachten nun die disjunkte Zerlegung $\mathcal{P}(A) = \mathcal{S} \cup \mathcal{T}$ mit

$$\mathcal{S} = \{B \in \mathcal{P}(A) \mid a \notin B\} \quad \text{und} \quad \mathcal{T} = \{B \in \mathcal{P}(A) \mid a \in B\}.$$

Nach Definition gilt $\mathcal{S} = \mathcal{P}(A')$, denn die Teilmengen von A' sind genau die Teilmengen $B \subseteq A$ mit $a \notin B$. Zwischen den Mengen $\mathcal{S}$ und $\mathcal{T}$ ist durch $\phi : \mathcal{S} \rightarrow \mathcal{T}$, $B \mapsto B \cup \{a\}$ eine Bijektion gegeben, denn $\psi : \mathcal{T} \rightarrow \mathcal{S}$, $B \mapsto B \setminus \{a\}$ ist offenbar die Umkehrabbildung von ϕ . Nach Proposition (4.13) folgt daraus $|\mathcal{S}| = |\mathcal{T}|$. Wir erhalten nun

$$|\mathcal{P}(A)| = |\mathcal{S}| + |\mathcal{T}| = 2|\mathcal{S}| = 2|\mathcal{P}(A')| = 2 \cdot 2^n = 2^{n+1}.$$

Damit ist der Induktionsschritt abgeschlossen. □



zum Induktionsschritt im Beweis von (4.7) (iv)

(4.16) Definition Für jede Menge B und jedes $k \in \mathbb{N}_0$ sei $\mathcal{P}_k(B)$ jeweils die Anzahl der k -elementigen Teilmengen von B , also

$$\mathcal{P}_k(B) = \left\{ A \in \mathcal{P}(B) \mid |A| = k \right\}.$$

Für alle $k, n \in \mathbb{N}_0$ definieren wir $\binom{n}{k} = |\mathcal{P}_k(M_n)|$ und bezeichnen diese Zahl als den **Binomialkoeffizienten** von n über k .

Beispielsweise ist $\binom{5}{3} = 10$, denn $M_5 = \{1, 2, 3, 4, 5\}$ hat genau zehn dreielementige Teilmengen, nämlich

$$\{1, 2, 3\}, \{1, 2, 4\}, \{1, 2, 5\}, \{1, 3, 4\}, \{1, 3, 5\}, \{1, 4, 5\}, \{2, 3, 4\}, \{2, 3, 5\}, \{2, 4, 5\}, \{3, 4, 5\}.$$

Einige Binomialkoeffizienten lassen sich direkt angeben. Für alle $k, n \in \mathbb{N}_0$ gilt

$$\binom{n}{0} = 1, \quad \binom{n}{1} = n \quad \text{und} \quad \binom{n}{k} = \binom{n}{n-k} \quad \text{falls } k \leq n, \quad \text{außerdem} \quad \binom{n}{k} = 0 \quad \text{falls } k > n.$$

Die Gleichung $\binom{n}{0} = 1$ ergibt sich aus der Feststellung, dass M_n nur eine nullelementige Teilmenge besitzt, nämlich die leere Menge. Für $n \in \mathbb{N}_0$ mit $n \geq 1$ sind die einelementigen Teilmengen von M_n offenbar genau die Mengen $\{1\}, \{2\}, \dots, \{n\}$, daraus folgt $\binom{n}{1} = n$. Sind $k, n \in \mathbb{N}_0$ mit $k > n$, dann gibt es nach Satz (4.15) (i) keine k -elementigen Teilmengen von M_n . Daraus folgt $\binom{n}{k} = 0$ für $k > n$.

Die Gleichung $\binom{n}{k} = \binom{n}{n-k}$ ist für $k \leq n$ ergibt sich durch folgende Überlegung: Ist $A \subseteq M_n$ eine k -elementige Teilmenge, dann ist $M_n \setminus A$ nach Satz (4.15) (ii) eine $n-k$ -elementige Teilmenge. Durch $\Phi : A \mapsto M_n \setminus A$ ist also eine Abbildung $\mathcal{P}_k(M_n) \rightarrow \mathcal{P}_{n-k}(M_n)$ gegeben. Diese besitzt $\mathcal{P}_{n-k}(M_n) \rightarrow \mathcal{P}_k(M_n)$, $B \mapsto M_n \setminus B$ als Umkehrabbildung. Denn wegen $n - (n-k) = k$ ist $M_n \setminus B$ für jedes $B \in \mathcal{P}_{n-k}(M_n)$ in $\mathcal{P}_k(M_n)$ enthalten, und wegen $M_n \setminus (M_n \setminus A) = A$, $M_n \setminus (M_n \setminus B) = B$ für alle $A \in \mathcal{P}_k(M_n)$ und $B \in \mathcal{P}_{n-k}(M_n)$ sind die beiden Abbildungen tatsächlich zueinander invers.

Die folgenden beiden Aussagen ermöglichen die Berechnung beliebiger Binomialkoeffizienten.

(4.17) Proposition Seien $k, n \in \mathbb{N}_0$, wobei $k \geq 1$ ist. Dann gilt

$$\binom{n+1}{k} = \binom{n}{k} + \binom{n}{k-1}.$$

Beweis: Wir betrachten die disjunkte Zerlegung von $\mathcal{P}_k(M_{n+1}) = \mathcal{S} \cup \mathcal{T}$ in die Teilmengen

$$\mathcal{S} = \left\{ A \in \mathcal{P}_k(M_{n+1}) \mid n+1 \in A \right\} \quad \text{und} \quad \mathcal{T} = \left\{ A \in \mathcal{P}_k(M_{n+1}) \mid n+1 \notin A \right\},$$

wobei $\mathcal{T} = \mathcal{P}_k(M_n)$ ist. Offenbar handelt es sich bei $\phi : \mathcal{P}_{k-1}(M_n) \rightarrow \mathcal{S}$, $A \mapsto A \cup \{n+1\}$ um eine Bijektion, denn $\psi : \mathcal{S} \rightarrow \mathcal{P}_{k-1}(M_n)$, $B \mapsto B \setminus \{n+1\}$ ist eine Umkehrabbildung von ϕ . Daraus folgt $|\mathcal{P}_{k-1}(M_n)| = |\mathcal{S}|$, und insgesamt erhalten wir $\binom{n+1}{k} = |\mathcal{P}_k(M_{n+1})| = |\mathcal{S}| + |\mathcal{T}| = |\mathcal{P}_{k-1}(M_n)| + |\mathcal{P}_k(M_n)| = \binom{n}{k-1} + \binom{n}{k}$. $\square$

Aus dieser Formel ergibt sich als Berechnungsschema das sogenannte *Pascalsche Dreieck*. Ist $\binom{n}{\ell}$ für jedes $\ell \in \mathbb{N}_0$ bereits bekannt, dann lässt sich der Wert $\binom{n+1}{k}$ dadurch berechnen, dass man die im Dreieck unmittelbar darüber stehenden Werte $\binom{n}{k-1}$ und $\binom{n}{k}$ einfach addiert.

$$\begin{array}{ccccccc}
 & & & & & & 1 \\
 & & & & & 1 & 1 \\
 & & \binom{1}{0} & & \binom{1}{1} & & \\
 & \binom{2}{0} & & \binom{2}{1} & & \binom{2}{2} & \\
 \binom{3}{0} & & \binom{3}{1} & & \binom{3}{2} & & \binom{3}{3}
 \end{array}
 \qquad
 \begin{array}{ccccccc}
 & & & & & & 1 \\
 & & & & 1 & & 1 \\
 & & 1 & & 2 & & 1 \\
 1 & & 3 & & 3 & & 1
 \end{array}$$

Um eine explizite, nicht-rekursive Formel für die Binomialkoeffizienten anzugeben, benötigen wir die sogenannte **Fakultätsfunktion**. Es handelt sich dabei um eine Abbildung $\mathbb{N}_0 \rightarrow \mathbb{N}$, $n \mapsto n!$, die rekursiv definiert ist durch

$$0! = 1 \quad \text{und} \quad (n+1)! = (n+1) \cdot n! \quad \text{für alle } n \in \mathbb{N}_0.$$

Es gilt also $1! = 1 \cdot 0! = 1$, $2! = 2 \cdot 1! = 2$, $3! = 3 \cdot 2! = 6$, $4! = 4 \cdot 3! = 24$, $5! = 5 \cdot 4! = 120$ usw. Insgesamt handelt es sich um eine sehr schnell wachsende Funktion. Zum Beispiel ist die Fakultät von 30 gegeben durch $30! = 265.252.859.812.191.058.636.308.480.000.000$, das sind immerhin schon 33 Dezimalstellen.

(4.18) Satz Für alle $k, n \in \mathbb{N}_0$ mit $k \leq n$ gilt

$$\binom{n}{k} = \frac{n!}{k!(n-k)!}.$$

Beweis: Wir beweisen die Aussage durch vollständige Induktion über $n \in \mathbb{N}_0$, jeweils für alle $k \in \mathbb{N}_0$ mit $k \leq n$. Für $n = k = 0$ folgt die Gleichung aus $\binom{0}{0} = 1 = \frac{0!}{0! \cdot 0!}$ erfüllt. Sei nun $n \in \mathbb{N}_0$ beliebig vorgegeben, und setzen wir die Gleichung für dieses n und alle $k \in \mathbb{N}_0$ mit $k \leq n$ voraus. Unter dieser Voraussetzung beweisen wir die Gleichung für $n+1$ und $1 \leq k \leq n+1$. Ist $k=0$, so gilt $\binom{n+1}{0} = 1 = \frac{(n+1)!}{0!(n+1)!}$. Für $1 \leq k \leq n$ folgt die Gleichung mit Hilfe von Proposition (4.17) aus der Rechnung

$$\begin{aligned}
 \binom{n+1}{k} &= \binom{n}{k} + \binom{n}{k-1} = \frac{n!}{k!(n-k)!} + \frac{n!}{(k-1)!(n-k+1)!} = \\
 &= \frac{n!(n-k+1)}{k!(n-k)!(n-k+1)} + \frac{n!k}{k(k-1)!(n-k+1)!} = \frac{n!((n-k+1)+k)}{k!(n-k+1)!} = \frac{(n+1)!}{k!(n+1-k)!}.
 \end{aligned}$$

Der einzige verbleibende Fall ist $k = n+1$. Weil M_{n+1} die einzige $(n+1)$ -elementige Menge von M_{n+1} ist, gilt $\binom{n+1}{n+1} = |\mathcal{P}_{n+1}(M_{n+1})| = 1$, und ebenso ist $\frac{(n+1)!}{(n+1)!0!}$ gleich 1. $\square$

Nachdem wir uns bis jetzt mit der Mächtigkeit *endlicher* Mengen beschäftigt haben, wenden wir uns nun den *unendlichen* Mengen zu. Durch Proposition (4.13) wird folgende Definition nahegelegt.

(4.19) Definition Wir bezeichnen zwei Mengen A und B als **gleichmächtig** (oder als Mengen mit gleicher Mächtigkeit) und schreiben $|A| = |B|$, wenn eine Bijektion $\phi : A \rightarrow B$ existiert.

Mit Hilfe dieser Definition können nun auch verschiedene Arten von unendlichen Mengen gegeneinander abgegrenzt werden.

(4.20) Definition Eine Menge A wird **abzählbar unendlich** genannt, wenn sie die gleiche Mächtigkeit wie $\mathbb{N}$ besitzt. Eine Menge, die endlich oder abzählbar unendlich ist, nennen wir **höchstens abzählbar**. Eine Menge, die nicht höchstens abzählbar ist, bezeichnet man als **überabzählbar**.

Als erstes bemerken wir

(4.21) Proposition Abzählbar unendliche Mengen sind unendlich.

Beweis: Nehmen wir an, dass eine Menge A zugleich endlich und abzählbar unendlich ist. Dann gibt es einerseits ein $n \in \mathbb{N}_0$ mit $|A| = n$, also eine Bijektion $\varphi : M_n \rightarrow A$, und andererseits eine Bijektion $\psi : \mathbb{N} \rightarrow A$. Sei ψ_{n+1} die Einschränkung von ψ auf M_{n+1} ; dann wäre $\varphi^{-1} \circ \psi_{n+1} : M_{n+1} \rightarrow M_n$ eine injektive Abbildung von M_{n+1} in eine echte Teilmenge von M_{n+1} . Aber eine solche Abbildung kann es nach Proposition (4.12) nicht geben. $\square$

Genau wie bei den endlichen Mengen sehen wir uns nun an, unter welchen Mengenoperationen die Abzählbarkeit erhalten bleibt.

(4.22) Lemma Jede unendliche Teilmenge von $\mathbb{N}$ ist abzählbar unendlich.

Beweis: Sei $A \subseteq \mathbb{N}$ eine unendliche Teilmenge. Wir definieren eine injektive Abbildung $\varphi : \mathbb{N} \rightarrow A$ durch folgende Rekursionsvorschrift: Wie in den Übungen mit dem Induktionsprinzip gezeigt wurde, besitzt jede nichtleere Teilmenge von $\mathbb{N}$ ein kleinstes Element. Bezeichnet a_1 das kleinste Element in A , dann setzen wir $\varphi(1) = a_1$. Sei nun $n \in \mathbb{N}$ und nehmen wir nun an, dass $\varphi(k)$ für $1 \leq k \leq n$ bereits definiert wurde, und dass φ auf $M_n = \{1, \dots, n\}$ injektiv ist.

Wäre die Menge $A \setminus \varphi(M_n)$ leer, dann würde daraus $\varphi(M_n) = A$ folgen. Damit wäre φ eine Bijektion zwischen M_n und A , was aber der Voraussetzung widerspricht, dass A unendlich ist. So aber können wir in $A \setminus \varphi(M_n)$ ein kleinstes Element b wählen und $\varphi(n+1) = b$ definieren. Offenbar ist φ auch auf M_{n+1} injektiv. Denn nehmen wir an, es gäbe Elemente $k, \ell \in M_{n+1}$ mit $k < \ell$ und $\varphi(k) = \varphi(\ell)$. Wegen der Injektivität von $\varphi|_{M_n}$ ist dies nur möglich, wenn $k \in M_n$ und $\ell = n+1$ ist. Aber dann ist $\varphi(\ell) = \varphi(n+1) = b \notin \varphi(M_n)$ und $\varphi(k) \in \varphi(M_n)$. Also können $\varphi(k)$ und $\varphi(\ell)$ nicht übereinstimmen.

Insgesamt erhalten wir so eine Abbildung $\varphi : \mathbb{N} \rightarrow A$. Auch diese ist injektiv. Sind nämlich $k, \ell \in \mathbb{N}$ mit $k < \ell$, dann folgt aus der Injektivität von $\varphi|_{M_\ell}$ sofort $\varphi(k) \neq \varphi(\ell)$. Es bleibt zu zeigen, dass φ auch surjektiv ist. Nehmen wir an, dies ist nicht der Fall. Dann existiert in A ein kleinstes Element a mit $a \notin \varphi(\mathbb{N})$. Seien $a_1, \dots, a_r$ die endlich vielen Elemente in A , die kleiner als a sind. Für jedes $i \in \{1, \dots, r\}$ gibt es ein $n_i \in \mathbb{N}$ mit $\varphi(n_i) = a_i$. Nach eventueller Vertauschung der Elemente $a_1, \dots, a_r$ können wir annehmen, dass n_r unter den Zahlen $a_1, \dots, a_r$ maximal ist. Daraus

folgt $\{a_1, \dots, a_r\} \subseteq \varphi(M_n)$ für $n = n_r$. Wir behaupten nun, dass $\varphi(n+1) = a$ gelten muss, im Widerspruch zur Annahme. Wegen $A \setminus \{a_1, \dots, a_r\} \supseteq A \setminus \varphi(M_n)$ ist a auch das kleinste Element in $A \setminus \varphi(M_n)$. Also ergibt sich die Gleichung $\varphi(n+1) = a$ direkt aus unserer Rekursionsvorschrift. $\square$

(4.23) Proposition

- (i) Teilmengen höchstens abzählbarer Mengen sind höchstens abzählbar.
- (ii) Sind A, B Mengen, ist A höchstens abzählbar und $\phi : A \rightarrow B$ eine surjektive Abbildung, dann ist auch B höchstens abzählbar.

Beweis: zu (i) Sei B eine höchstens abzählbare Menge und $A \subseteq B$. Ist B endlich, dann ist A nach Satz (4.15) (i) ebenfalls endlich und damit höchstens abzählbar. Wir können deshalb annehmen, dass B abzählbar unendlich ist. Demnach gibt es eine Bijektion $\varphi : B \rightarrow \mathbb{N}$. Ist A endlich, dann ist A nach Definition höchstens abzählbar. Wir können also annehmen, dass A unendlich ist. Auf Grund der Bijektivität von φ ist $\varphi(A)$ eine unendliche Teilmenge von $\mathbb{N}$. Diese ist nach Lemma (4.22) abzählbar unendlich. Also ist auch A abzählbar unendlich, insbesondere höchstens abzählbar.

zu (ii) Nach Satz (4.8) (ii) existiert eine Abbildung $\psi : B \rightarrow A$ mit $\varphi \circ \psi = \text{id}_B$, und nach Satz (4.8) (i) ist diese Abbildung injektiv. Wie unter (i) gezeigt, ist $\psi(B)$ als Teilmenge der höchstens abzählbaren Menge A ebenfalls höchstens abzählbar. Weil ψ als injektive Abbildung eine zwischen B und $\psi(B)$ bijektiv ist, ist auch B höchstens abzählbar. $\square$

Abbildungen können genutzt werden, um Elemente einer Menge A durch Elemente einer anderen Menge I zu indizieren. In diesem Zusammenhang bezeichnet man eine Abbildung $\varphi : I \rightarrow A$ auch als **Familie** von Elementen der Menge A und I als **Indexmenge** der Familie. Man verwendet dann für die Abbildung φ die Notation $(a_i)_{i \in I}$, und das Element $\varphi(i) \in A$ bezeichnet man mit a_i . Ist $I = \mathbb{N}$ oder $\mathbb{N}_0$, dann nennt man die Familie auch eine **Folge**.

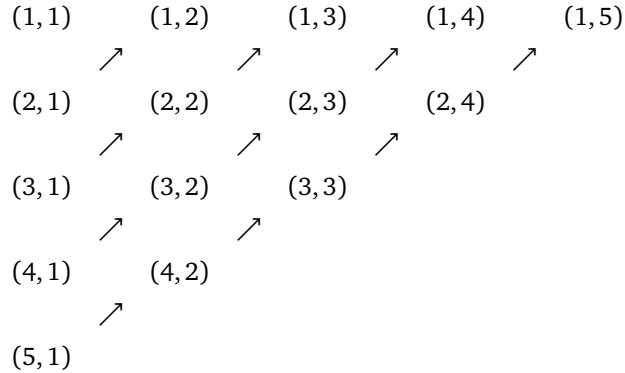
Beispielsweise wird durch $a_1 = 3, a_2 = 5, a_3 = 97, a_4 = 3$ eine Familie $(a_i)_{i \in I}$ natürlicher Zahlen mit der Indexmenge $\{1, 2, 3, 4\}$ definiert. Durch die Festlegung $a_n = n^2$ für alle $n \in \mathbb{N}$ erhält man eine Folge $(a_n)_{n \in \mathbb{N}}$ natürlicher Zahlen, nämlich die Folge der Quadratzahlen.

Familien werden verwendet, um auf die Elemente einer Menge A leichter zugreifen zu können (ähnlich wie die Seitennummern einen leichteren Zugriff auf die Seiten eines Buchs ermöglichen). Man kann auf diese Weise auch bestimmte Elemente von A auszeichnen oder sie (im Fall von $I = \mathbb{N}$) in eine bestimmte Reihenfolge bringen. Zu beachten ist dabei, dass die Abbildung $I \rightarrow A, i \mapsto a_i$ im allgemeinen weder injektiv noch surjektiv zu sein braucht. Beispielsweise kann dasselbe Element von A in einer Familie auch mehrfach vorkommen, also $a_i = a_j$ für verschiedene $i, j \in I$ gelten.

(4.24) Satz

- (i) Sind A und B abzählbar unendliche Mengen, dann ist auch $A \times B$ abzählbar unendlich.
- (ii) Ist I höchstens abzählbar, und ist $(A_i)_{i \in I}$ eine Familie bestehend aus lauter höchstens abzählbaren Mengen A_i , dann ist auch die Vereinigung $\bigcup_{i \in I} A_i$ höchstens abzählbar.

Beweis: zu (i) Zunächst führen wir den Beweis auf den Fall $A = B = \mathbb{N}$ zurück. Weil A und B abzählbar unendlich sind, gibt es Bijektionen $\varphi : \mathbb{N} \rightarrow A$ und $\psi : \mathbb{N} \rightarrow B$. Man überprüft leicht, dass dann die Abbildung $\mathbb{N} \times \mathbb{N} \rightarrow A \times B$, $(m, n) \mapsto (\varphi(m), \psi(n))$ ebenfalls bijektiv ist. Wenn also $\mathbb{N} \times \mathbb{N}$ abzählbar unendlich ist, dann gilt dasselbe für $A \times B$. Es genügt also nachzuweisen, dass $\mathbb{N} \times \mathbb{N}$ abzählbar unendlich ist. Dazu geben wir eine injektive Abbildung zwischen $\mathbb{N} \times \mathbb{N}$ und $\mathbb{N}$ an. Die grundlegende Idee besteht darin, die Paare in $\mathbb{N} \times \mathbb{N}$ nach dem folgenden Schema durchnummerieren.



Dies wird realisiert durch die Abbildung $\phi : \mathbb{N}^2 \rightarrow \mathbb{N}$, $(m, n) \mapsto n + \frac{1}{2}(m + n - 2)(m + n - 1)$, die folgendermaßen zu Stande kommt: In jeder Diagonale des angegebenen Schemas befinden sich die Paare $(m, n) \in \mathbb{N}^2$ mit konstanter Summe $m + n$, wobei (m, n) jeweils auf der n -ten Position der $(m + n - 1)$ -ten Diagonale landet. Die r -te Diagonale hat für jedes $r \in \mathbb{N}$ jeweils genau r Einträge. Die Formel $1 + \dots + r = \frac{1}{2}r(r + 1)$, die in Satz (2.4) bewiesen wurde, zeigt, dass die $(m + n - 2)$ vorausgehenden Diagonalen insgesamt die Länge $\frac{1}{2}(m + n - 2)(m + n - 1)$ haben. Also landet das Element (m, n) im Schema auf der Position $\frac{1}{2}(m + n - 2)(m + n - 1) + n$.

Wir beweisen nun die Injektivität der Abbildung ϕ . Dazu seien $(m_1, n_1), (m_2, n_2) \in \mathbb{N}^2$ mit $\phi(m_1, n_1) = \phi(m_2, n_2)$ vorgegeben. Zu zeigen ist $(m_1, n_1) = (m_2, n_2)$. Als erstes überprüfen wir, dass die Zahlen

$$r_1 = m_1 + n_1 - 2 \quad \text{und} \quad r_2 = m_2 + n_2 - 2$$

übereinstimmen. Nehmen wir an, dass dies nicht der Fall ist und zum Beispiel $r_1 < r_2$ gilt. Dann ist $\phi(m_1, n_1) \leq \frac{1}{2}r_1(r_1 + 1) + r_1$ und $\phi(m_2, n_2) \geq \frac{1}{2}r_2(r_2 + 1)$. Wir erhalten

$$\begin{aligned}
 \phi(m_2, n_2) - \phi(m_1, n_1) &\geq \frac{1}{2}r_2(r_2 + 1) - \left(\frac{1}{2}r_1(r_1 + 1) + r_1\right) \geq \frac{1}{2}(r_1 + 1)(r_1 + 2) - \left(\frac{1}{2}r_1(r_1 + 1) + r_1\right) \\
 &= \frac{1}{2}r_1^2 + \frac{3}{2}r_1 + 1 - \left(\frac{1}{2}r_1^2 + \frac{3}{2}r_1\right) = 1,
 \end{aligned}$$

was der Voraussetzung $\phi(m_1, n_1) = \phi(m_2, n_2)$ widerspricht. Also muss $r_1 = r_2$ gelten. Aus $r_1 = r_2$ folgt aber direkt $\frac{1}{2}r_1(r_1 + 1) = \frac{1}{2}r_2(r_2 + 1)$. Zusammen mit $\phi(m_1, n_1) = \phi(m_2, n_2) \Leftrightarrow \frac{1}{2}r_1(r_1 + 1) + n_1 = \frac{1}{2}r_2(r_2 + 1) + n_2$ folgt daraus $n_1 = n_2$ und damit auch $m_1 = m_2$. Damit ist die Injektivität bewiesen.

Wir haben somit gezeigt, dass $\mathbb{N}^2$ gleichmächtig zur Teilmenge $\psi(\mathbb{N}^2)$ von $\mathbb{N}$ ist. Nach Lemma (4.22) ist $\mathbb{N}^2$ höchstens abzählbar. Andererseits ist $\mathbb{N}_0 \times \mathbb{N}_0$ unendlich; wäre dies nicht so, dann würde aus der Injektivität der Abbildung $\mathbb{N} \rightarrow \mathbb{N}^2$, $m \mapsto (m, 1)$ auch die Endlichkeit von $\mathbb{N}$ folgen, was nach Proposition (4.14) ausgeschlossen ist.

zu (ii) Da I höchstens abzählbar ist, gibt es eine injektive Abbildung $\varphi : I \rightarrow \mathbb{N}$. Auf Grund der Injektivität gibt es nach Satz (4.8) (i) eine Abbildung $\psi : \mathbb{N} \rightarrow I$ mit $\psi \circ \varphi = \text{id}_I$. Diese ist nach Satz (4.8) (ii) surjektiv. Dasselbe Argument liefert uns für jedes $i \in I$ auch eine surjektive Abbildung $\varphi_i : \mathbb{N} \rightarrow A_i$. Wir behaupten nun, dass durch

$$\phi : \mathbb{N}^2 \longrightarrow \bigcup_{i \in I} A_i, \quad (m, n) \mapsto \varphi_{\psi(m)}(n)$$

ebenfalls eine surjektive Abbildung gegeben ist. Ist nämlich $a \in \bigcup_{i \in I} A_i$ vorgegeben, dann gibt es $i \in I$ mit $a \in A_i$, ein $m \in \mathbb{N}_0$ mit $\psi(m) = i$ und ein $n \in \mathbb{N}_0$ mit $\varphi_i(n) = a$. Es folgt $\phi(m, n) = \varphi_{\psi(m)}(n) = \varphi_i(n) = a$. Nach Proposition (4.23) (ii) ist deshalb mit $\mathbb{N}^2$ auch die Menge $\bigcup_{i \in I} A_i$ höchstens abzählbar. $\square$

§ 5. Algebraische Grundstrukturen und Matrizen

Inhaltsübersicht

In diesem Abschnitt definieren wir eine Reihe grundlegender algebraischer Strukturen, die im weiteren Verlauf der Vorlesung eine wichtige Rolle spielen werden: Halbgruppen, Monoide, Gruppen, Ringe und Körper. Ein konkretes Beispiel für einen Ring sind die ganzen Zahlen $\mathbb{Z}$. Bei den rationalen und reellen Zahlen, $\mathbb{Q}$ und $\mathbb{R}$, handelt es sich sogar um Körper. Mit Hilfe der in § 3 eingeführten Kongruenzrelationen und Kongruenzklassen werden wir sehen, dass es auch Ringe und Körper mit nur endlich vielen Elementen gibt.

Um weitere Beispiele für solche Strukturen zu erhalten (und weil wir sie als wichtiges Hilfsmittel der Linearen Algebra brauchen werden), führen wir in der zweiten Hälfte des Kapitels die Matrizen über einem beliebigen Ring ein, zusammen mit entsprechend angepassten Rechenoperationen, die man wie bei den Zahlen als „Addition“ und „Multiplikation“ bezeichnet. Die Matrizen unterscheiden sich von zuvor behandelten Strukturen unter anderem dadurch, dass die Multiplikation auf ihnen nicht mehr kommutativ ist, also in der Regel $AB \neq BA$ gilt.

Wichtige Begriffe und Sätze

- Definition der Verknüpfungen auf einer Menge, Eigenschaften „assoziativ“ und „kommutativ“
- Abgeschlossenheit einer Teilmenge unter einer Verknüpfung
- Halbgruppen, Monoide und Gruppen
- Neutralelement in einer Halbgruppe
- invertierbares Element in einem Monoid
- Ringe und Körper
- Restklassenringe und Restklassenkörper
- Matrizen, Nullmatrix, Einheitsmatrix, invertierbare Matrix

Unsere Einführung in das Thema dieses Kapitels beginnt mit dem Begriff der Verknüpfung.

(5.1) Definition Eine **Verknüpfung** auf einer Menge A ist eine Abbildung $A \times A \rightarrow A$.

Beispiele für Verknüpfungen sind die Addition, die Subtraktion und die Multiplikation auf der Menge $\mathbb{Z}$ der ganzen Zahlen, der Menge $\mathbb{Q}$ der rationalen Zahlen oder der Menge $\mathbb{R}$ der reellen Zahlen. Die Division ist *keine* Verknüpfung auf $\mathbb{Z}$, $\mathbb{Q}$ oder $\mathbb{R}$, weil beispielsweise die Division durch 0 unzulässig ist. Um eine Verknüpfung zu erhalten, müsste man die Division geeignet einschränken. Sie wird zum Beispiel zu einer Verknüpfung auf der Teilmenge $\mathbb{Q}^\times = \mathbb{Q} \setminus \{0\}$. Ebenso sind Addition und Multiplikation Verknüpfungen auf $\mathbb{N}$ und $\mathbb{N}_0$, wohingegen Subtraktion und Division keine Verknüpfungen auf diesen Mengen sind. Beispielsweise liefert die Subtraktion zwar eine Abbildung $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{Z}$, aber keine Abbildung $\mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$.

Als Bezeichnungen für eine Verknüpfung sind die Symbole $\cdot$, $\odot$, $*$, $+$, $\oplus$ und einige Varianten üblich. Wird eines der Symbole $\cdot$, $\odot$, $*$ verwendet, dann spricht man von einer **multiplikativen** Verknüpfung, bei $+$ oder $\oplus$ nennt man sie **additiv**. Die beiden Typen unterscheiden sich aber ausschließlich durch das verwendete Symbol, mathematisch gesehen besteht zwischen einer additiven und einer multiplikativen Verknüpfung keinerlei Unterschied.

Multiplikative Verknüpfungssymbole werden zur Vereinfachung der Notation häufig auch weggelassen, d.h. an Stelle von $a \cdot b$ schreibt man einfach ab . Sollen mehrere Elemente miteinander verknüpft werden, so ist die Verwendung von **Klammern** üblich, um die Reihenfolge der angewendeten Verknüpfungen anzuzeigen. So bedeutet zum Beispiel der Ausdruck $a(b(cd))$, dass zunächst das Element $x_1 = cd$ gebildet wird, anschließend $x_2 = bx_1$ und schließlich $x_3 = ax_2$.

(5.2) Definition Eine Verknüpfung $\cdot$ auf einer Menge A bezeichnet man als

- (i) **kommutativ**, wenn $ab = ba$ für alle $a, b \in A$
- (ii) **assoziativ**, wenn $a(bc) = (ab)c$ für alle $a, b, c \in A$ erfüllt ist.

Man bezeichnet eine Teilmenge $B \subseteq A$ als **abgeschlossen** unter der Verknüpfung $\cdot$, wenn für alle $b, b' \in B$ jeweils $bb' \in B$ gilt. Man erhält in diesem Fall eine Verknüpfung $\cdot_B$ auf B , indem man $b \cdot_B b' = bb'$ für alle $b, b' \in B$ setzt.

Bei assoziativen Verknüpfungen können die Klammern auch weggelassen werden. Für beliebige Elemente $a, b, c, d \in A$ ist dann zum Beispiel $abcd$ eine Kurzschreibweise für das Element $a(b(cd))$, welches auf Grund der Assoziativität mit jedem anders geklammerten Ausdruck, etwa $(ab)(cd)$ oder $a((bc)d)$, übereinstimmt.

Viele der bekannten Verknüpfungen sind assoziativ und kommutativ, beispielsweise die Addition und die Multiplikation auf den Zahlbereichen $\mathbb{N}$, $\mathbb{N}_0$, $\mathbb{Z}$, $\mathbb{Q}$ und $\mathbb{R}$. Die Subtraktion auf $\mathbb{Z}$ ist aber weder kommutativ noch assoziativ. Beispielsweise ist $2 - 1 = 1 \neq -1 = 1 - 2$, und $1 - (1 - 1) = 1 - 0 = 1 \neq -1 = 0 - 1 = (1 - 1) - 1$.

Auch Beispiele für abgeschlossene Teilmengen bezüglich der bekannten Verknüpfungen lassen sich in großer Zahl finden. Betrachtet man beispielsweise die Kette der Inklusionen $\mathbb{N} \subseteq \mathbb{N}_0 \subseteq \mathbb{Z} \subseteq \mathbb{Q} \subseteq \mathbb{R}$, dann ist jede Teilmenge in dieser Kette bezüglich Addition und Multiplikation abgeschlossen in ihrem Nachfolger. Auch die Teilmenge $2\mathbb{Z}$ in $\mathbb{Z}$ ist abgeschlossen bezüglich Addition und Multiplikation, ebenso die einelementige Teilmenge $\{0\}$. Die Menge $\{1\}$ ist ebenfalls abgeschlossen bezüglich Multiplikation, aber nicht bezüglich Addition, denn es gilt $1 + 1 = 2 \notin \{1\}$.

Nach diesen Vorbereitungen können wir nun die ersten algebraischen Grundstrukturen definieren.

(5.3) Definition

- (i) Eine **Halbgruppe** ist ein Paar $(G, \cdot)$ bestehend aus einer nichtleeren Menge G und einer assoziativen Verknüpfung $\cdot$ auf G .
- (ii) Ein Element $e \in G$ in einer Halbgruppe wird **Neutralelement** genannt, wenn $g \cdot e = e \cdot g = g$ für alle $g \in G$ erfüllt ist.
- (iii) Eine Halbgruppe $(G, \cdot)$, in der ein Neutralelement existiert, wird **Monoid** genannt.

Darauf aufbauend definieren wir weiter

(5.4) Definition

- (i) Ein Element g in einem Monoid $(G, \cdot)$ mit Neutralelement e heißt **invertierbar**, wenn ein $h \in G$ existiert, so dass die Gleichungen $g \cdot h = h \cdot g = e$ erfüllt sind.
- (ii) Eine **Gruppe** ist ein Monoid, in dem jedes Element invertierbar ist.
- (iii) Eine Halbgruppe, und ebenso ein Monoid und eine Gruppe, wird **kommutativ** oder **abelsch** genannt, wenn die zugehörige Verknüpfung kommutativ ist.

Auch für diese neuen Begriffe liefern die bekannten Zahlbereiche eine Vielzahl von Beispielen.

- (i) Das Paar $(\mathbb{N}, +)$ ist eine Halbgruppe, aber kein Monoid; es existiert kein Neutralelement, weil für alle $a, b \in \mathbb{N}$ stets $a + b \neq a$ gilt. Das Paar $(\mathbb{N}, \cdot)$ ist ein Monoid, mit 1 als Neutralelement, denn es gilt $1 \cdot a = a \cdot 1 = a$ für alle $a \in \mathbb{N}$. Das Paar ist aber keine Gruppe, denn es gibt beispielsweise kein $a \in \mathbb{N}$ mit $2 \cdot a = 1$ (und wie wir gleich sehen werden, existiert in jedem Monoid immer nur ein Neutralelement).
- (ii) Die Paare $(\mathbb{N}_0, +)$ ist ein Monoid, mit 0 als Neutralelement. Ebenso ist $(\mathbb{N}_0, \cdot)$ ein Monoid, das Neutralelement ist hier aber die 1. Beide Strukturen sind aber keine Gruppen.
- (iii) Das Paar $(\mathbb{Z}, +)$ ist sogar eine Gruppe, denn die 0 ist Neutralelement, und für jedes $a \in \mathbb{Z}$ gilt $a + (-a) = (-a) + a = 0$. Das Paar $(\mathbb{Z}, \cdot)$ ist zwar ein Monoid, aber keine Gruppe.
- (iv) Das Paar $(\mathbb{Q}, +)$ ist eine Gruppe, während $(\mathbb{Q}, \cdot)$ ein Monoid, aber keine Gruppe ist. Schränkt man die Verknüpfung $\cdot$ allerdings auf $\mathbb{Q}^\times$ ein, so erhält man eine Gruppe. Denn für jedes Element $a \in \mathbb{Q}^\times$ können wir den Kehrwert a^{-1} bilden, und es gilt $a \cdot a^{-1} = a^{-1} \cdot a = 1$.
- (v) Beim Körper $\mathbb{R}$ der reellen Zahlen haben wir dieselben Verhältnisse: Das Paar $(\mathbb{R}, +)$ ist eine Gruppe, $(\mathbb{R}, \cdot)$ ist ein Monoid, aber keine Gruppe, und durch Einschränkung der Multiplikation auf $\mathbb{R}^\times = \mathbb{R} \setminus \{0\}$ erhält man auch hier eine Gruppe.

Alle hier aufgezählten Halbgruppen, Monoide und Gruppen sind abelsch, weil die Addition und die Multiplikation auf allen Zahlbereichen kommutative Verknüpfungen sind.

(5.5) Lemma Sei $(G, \cdot)$ ein Monoid, und sei e ein Neutralelement des Monoids. Dann gilt:

- (i) Das Monoid besitzt keine weiteren Neutralelemente. Man bezeichnet e deshalb als das Neutralelement des Monoids, und bezeichnet es mit e_G .
- (ii) Zu jedem invertierbaren Element $a \in G$ gibt es genau ein Element $b \in G$ mit $ab = ba = e_G$. Man nennt b das zu a **inverse Element**, und bezeichnet es mit a^{-1} .
- (iii) Ist a invertierbar und $b \in G$ ein Element mit $ab = e_G$, dann folgt daraus bereits $b = a^{-1}$. Ebenso folgt bereits aus der Gleichung $ba = e_G$, dass b das Inverse von a ist.
- (iv) Das Neutralelement ist invertierbar, und es gilt $e_G^{-1} = e_G$.
- (v) Sind a und b invertierbare Elemente des Monoids, dann sind auch ab und a^{-1} invertierbar, und es gilt $(ab)^{-1} = b^{-1}a^{-1}$ und $(a^{-1})^{-1} = a$.

Beweis: zu (i) Angenommen, e' ist ein weiteres Neutralelement des Monoids. Weil e ein Neutralelement ist, gilt $ee' = e'$. Weil e' Neutralelement ist, gilt auch $ee' = e$. Insgesamt gilt also $e = ee' = e'$.

zu (ii) Sei $a \in G$ invertierbar, und seien $b, c \in G$ Elemente mit $ab = ba = e_G$ und $ac = ca = e_G$. Dann gilt insgesamt $b = be_G = b(ac) = (ba)c = e_Gc = c$.

zu (iii) Sei a invertierbar, und sei $b \in G$ mit $ab = e_G$. Dann gilt $a^{-1} = a^{-1}e_G = a^{-1}(ab) = (a^{-1}a)b = e_Gb = b$. Setzen wir $ba = e_G$ voraus, dann erhalten wir ebenso $a^{-1} = e_Ga^{-1} = (ba)a^{-1} = b(aa^{-1}) = be_G = b$.

zu (iv) Die Gleichung $e_Ge_G = e_G$ zeigt, dass das Element e_G sein eigenes Inverses ist, also $e_G^{-1} = e_G$ gilt.

zu (v) Die Gleichungen $(ab)(b^{-1}a^{-1}) = a(bb^{-1})a^{-1} = (ae_G)a^{-1} = aa^{-1} = e_G$ und $(b^{-1}a^{-1})(ab) = b^{-1}(a^{-1}a)b = b^{-1}(e_Gb) = b^{-1}b = e_G$ zeigen, dass $b^{-1}a^{-1}$ das Inverse von ab ist, also insbesondere $(ab)^{-1} = b^{-1}a^{-1}$ gilt. Ebenso zeigen die Gleichungen $aa^{-1} = e_G$ und $a^{-1}a = e_G$, dass a das Inverse von a^{-1} ist, also $(a^{-1})^{-1} = a$ gilt. $\square$

(5.6) Folgerung Sei $(G, \cdot)$ ein Monoid, und sei $G^\times \subseteq G$ die Teilmenge der invertierbaren Elemente. Dann ist $G^\times$ bezüglich $\cdot$ abgeschlossen, und $(G^\times, \cdot_{G^\times})$ ist eine Gruppe.

Beweis: Die Abgeschlossenheit von $G^\times$ bezüglich $\cdot$ folgt direkt aus Teil (v) von Lemma (5.5). Somit können wir auf $G^\times$ durch $a \cdot_{G^\times} b = ab$ für alle $a, b \in G^\times$ eine Verknüpfung definieren; der einzige Unterschied zwischen den Abbildungen $\cdot$ und $\cdot_{G^\times}$ ist demnach der Definitions- und der Wertebereich. Die neue Verknüpfung $\cdot_{G^\times}$ ist assoziativ, denn für alle $a, b, c \in G^\times$ gilt

$$a \cdot_{G^\times} (b \cdot_{G^\times} c) = a(bc) = (ab)c = (a \cdot_{G^\times} b) \cdot_{G^\times} c.$$

Somit ist $(G^\times, \cdot_{G^\times})$ eine Halbgruppe. Nach Teil (iv) von Lemma (5.5) ist e_G in $G^\times$ enthalten. Außerdem gilt $a \cdot_{G^\times} e_G = ae_G = a$ und $e_G \cdot_{G^\times} a = e_Ga = a$ für alle $a \in G^\times$. Dies zeigt, dass $(G^\times, \cdot_{G^\times})$ ein Monoid ist, mit e_G als Neutralelement. Für jedes $a \in G^\times$ liegt auf Grund des Lemmas auch a^{-1} in $G^\times$, und es gilt $a \cdot_{G^\times} a^{-1} = aa^{-1} = e_G$ und $a^{-1} \cdot_{G^\times} a = a^{-1}a = e_G$.

Das Element a ist also im Monoid $(G^\times, \cdot_{G^\times})$ invertierbar. Weil $a \in G^\times$ beliebig vorgegeben war, folgt daraus, dass jedes Element in dem Monoid invertierbar und $(G^\times, \cdot_{G^\times})$ somit eine Gruppe ist. $\square$

Die Folgerung zeigt noch einmal, dass $\mathbb{Q}^\times$ und $\mathbb{R}^\times$ mit der Multiplikation jeweils eine Gruppe bilden, denn $(\mathbb{Q}, \cdot)$ und $(\mathbb{R}, \cdot)$ sind Monoide, und $\mathbb{Q}^\times$ bzw. $\mathbb{R}^\times$ sind genau die invertierbaren Elemente des Monoids. Im Monoid $(\mathbb{Z}, \cdot)$ gibt es nur zwei invertierbare Elemente, nämlich ± 1 . In diesem Fall zeigt die Folgerung, dass durch $(\{\pm 1\}, \cdot)$ eine zweielementige Gruppe gegeben ist.

Monoide und Gruppen lassen sich zu komplexeren algebraischen Strukturen kombinieren.

(5.7) Definition Ein **Ring** ist ein Tripel $(R, +, \cdot)$ bestehend aus einer Menge R und zwei Verknüpfungen $+$ und $\cdot$ auf R mit folgenden Eigenschaften.

- (i) Das Paar $(R, +)$ ist eine abelsche Gruppe.
- (ii) Das Paar $(R, \cdot)$ ist ein abelsches Monoid.
- (iii) Es gilt das Distributivgesetz $a(b + c) = ab + ac$ für alle $a, b, c \in R$.

Das Neutralelement der Gruppe $(R, +)$ wird das **Nullelement** des Rings genannt und mit 0_R bezeichnet. Das Neutralelement des Monoids $(R, \cdot)$ nennt man das **Einselement** des Rings und bezeichnet es mit 1_R . Wenn die Menge $R^\times$ der invertierbaren Elemente des Monoids $(R, \cdot)$ mit $R \setminus \{0_R\}$ übereinstimmt, dann nennt man $(R, +, \cdot)$ auch einen **Körper**.

Für jedes Element a in einem Ring R bezeichnen wir das Inverse von a in der Gruppe $(R, +)$ mit $-a$ und nennen es das **Negative** von a . Es gilt also $a + (-a) = (-a) + a = 0_R$ für alle $a \in R$. Ist a im Monoid $(R, \cdot)$ ein invertierbares Element, so bezeichnen wir das Inverse mit a^{-1} und nennen es den **Kehrwert** von a . Nach Definition gilt $a \cdot a^{-1} = a^{-1} \cdot a = 1_R$ für dieses Element a . Auch für die Ringe stellen wir einige Rechenregeln zusammen.

(5.8) Proposition Sei $(R, +, \cdot)$ ein Ring, und seien $a, b \in R$. Dann gilt

- (i) $-0_R = 0_R$, $-(-a) = a$ und $-(a + b) = (-a) + (-b)$,
- (ii) $0_R \cdot a = 0_R$, $a(-b) = (-a)b = -(ab)$ und $ab = (-a)(-b)$,
- (iii) $1_R^{-1} = 1_R$, $(a^{-1})^{-1} = a$ und $(ab)^{-1} = a^{-1}b^{-1}$, sofern a und b in $(R, \cdot)$ invertierbar sind.
- (iv) Ist R sogar ein Körper, dann folgt aus $ab = 0_R$ immer $a = 0_R$ oder $b = 0_R$.

Beweis: Die Regeln (i) und (iii) erhält man unmittelbar durch Anwendung von Lemma (5.5) auf die Gruppe $(R, +)$ bzw. das Monoid $(R, \cdot)$. Zum Beweis der Regel (ii) seien $a, b \in R$ vorgegeben. Die erste Gleichung erhält man durch die Rechnung

$$\begin{aligned} 0_R \cdot a &= 0_R \cdot a + 0_R &= 0_R \cdot a + 0_R \cdot a + (-0_R \cdot a) &= (0_R + 0_R) \cdot a + (-0_R \cdot a) \\ &= 0_R \cdot a + (-0_R \cdot a) &= 0_R. \end{aligned}$$

Die Gleichung $ab + a(-b) = a(b + (-b)) = a \cdot 0_R = 0_R$ zeigt, dass $a(-b)$ das Inverse von ab in der Gruppe $(R, +)$ ist, also $a(-b) = -(ab)$ gilt. Ebenso erhält man die Gleichung $(-a)b = -(ab)$. Die letzte Gleichung unter (ii) erhält man schließlich mit Hilfe der bereits hergeleiteten Regeln durch die Rechnung $(-a)(-b) = -(a(-b)) = -(-(ab)) = ab$.

Für den Beweis von (iv) setzen wir voraus, dass R ein Körper ist. Seien $a, b \in R$ mit $ab = 0_R$ vorgegeben, und nehmen wir an, dass a ungleich 0_R ist. Dann ist a auf Grund der Körpereigenschaft ein invertierbares Element bezüglich der Multiplikation. Es folgt dann $b = 1_R \cdot b = (a^{-1}a) \cdot b = a^{-1}(ab) = a^{-1} \cdot 0_R = 0_R$. $\square$

Allgemein ist es üblich, den Ausdruck $a + (-b)$ für $a, b \in R$ durch $a - b$ abzukürzen. Auf diese Weise erhält man eine neue Verknüpfung $-$ auf R . Wieder untersuchen wir die bekannten Zahlbereiche im Hinblick auf die soeben eingeführten Begriffe.

- (i) Die Mengen $\mathbb{N}$ und $\mathbb{N}_0$ bilden mit der Addition und der Multiplikation keine Ringe, denn $(\mathbb{N}, +)$ und $(\mathbb{N}_0, +)$ sind keine Gruppen.
- (ii) Die Menge $\mathbb{Z}$ der ganzen Zahlen bildet mit der Addition und der Multiplikation einen Ring. Es handelt sich aber um keinen Körper, denn wie wir oben festgestellt haben, ist die Menge der invertierbaren Elemente im Monoid $(\mathbb{Z}, \cdot)$ gleich $\{\pm 1\}$, sie stimmt also nicht mit $\mathbb{Z} \setminus \{0\}$ überein.
- (iii) Die rationalen Zahlen $\mathbb{Q}$ und die reellen Zahlen $\mathbb{R}$ bilden mit der Addition und der Multiplikation Körper, und somit auch Ringe.

In einem Ring R kann es durchaus passieren, dass Eins- und Nullelement zusammenfallen, also $0_R = 1_R$ gilt. Dies ist aber nur möglich, wenn der Ring nur aus einem einzigen Element besteht, also $R = \{0_R\} = \{1_R\}$ gilt. Setzen wir nämlich $0_R = 1_R$ voraus und ist $a \in R$ ein beliebiges Element, dann folgt $a = 1_R \cdot a = 0_R \cdot a = 0_R$, nach Teil (ii) von Proposition (5.8). In einem Körper K gilt dagegen immer $0_K \neq 1_K$, weil Ringe R bestehend aus nur einem Element nach Definition keine Körper sind: Hier ist die Menge $R^\times$ der invertierbaren Elemente gleich $\{0_R\}$, während $R \setminus \{0_R\} = \emptyset$ ist. Die Gleichung $R^\times = R \setminus \{0_R\}$ ist hier also nicht erfüllt.

Um die Definition der Ringe noch besser illustrieren, führen wir eine weitere Klassen von Beispielen ein, die im Gegensatz zu den bisherigen nicht durch die aus der Schulmathematik bekannten Zahlbereiche zu Stande kommt.

Wir haben in § 3 die Mengen $\mathbb{Z}/n\mathbb{Z}$ bestehend aus den Kongruenzklassen $\bar{k} = [k]_n = k + n\mathbb{Z}$ mit $k \in \mathbb{Z}$ eingeführt. Dort haben wir auch gesehen, dass auf Grund der Äquivalenz $\bar{k} = \bar{\ell} \Leftrightarrow k \equiv_n \ell \Leftrightarrow n \mid (k - \ell)$ diese Menge aus n verschiedenen Elementen besteht, die durch $\bar{0}, \bar{1}, \dots, \overline{n-1}$ angegeben werden können. In der Algebra-Vorlesung werden wir zeigen, dass man auch auf $\mathbb{Z}/n\mathbb{Z}$ auf natürliche Weise eine Addition und eine Multiplikation definieren kann.

(5.9) Satz Sei $n \in \mathbb{N}$. Dann gibt es eindeutig bestimmte Abbildungen $+: \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ und $\cdot: \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ mit

$$(a + n\mathbb{Z}) + (b + n\mathbb{Z}) = (a + b) + n\mathbb{Z} \quad \text{und} \quad (a + n\mathbb{Z}) \cdot (b + n\mathbb{Z}) = ab + n\mathbb{Z}$$

für alle $a, b \in \mathbb{Z}$.

Vollständig lassen sich diese Verknüpfungen durch **Verknüpfungstabellen** beschreiben, in denen jede Summe bzw. jedes Produkt von je zwei Elementen der Menge $\mathbb{Z}/n\mathbb{Z}$ angegeben ist. Für die Angabe wird dabei einheitlich die Darstellung der Elemente von $\mathbb{Z}/n\mathbb{Z}$ durch das Repräsentantensystem $\{0, 1, \dots, n-1\}$ der Äquivalenzklassen verwendet. In $\mathbb{Z}/7\mathbb{Z}$ gilt beispielsweise $\bar{2} + \bar{2} = \bar{4}$, $\bar{3} + \bar{5} = \bar{8} = \bar{1}$ und $\bar{5} + \bar{6} = \bar{11} = \bar{4}$. Dabei wurde jeweils im ersten Schritt die Definition der Verknüpfung verwendet, und im zweiten Schritt wurde die Darstellung des Elements auf das angegebene Repräsentantensystem umgerechnet. Nach demselben Schema erhält man $\bar{2} \cdot \bar{3} = \bar{6}$, $\bar{3} \cdot \bar{5} = \bar{15} = \bar{1}$ und $\bar{5} \cdot \bar{6} = \bar{30} = \bar{2}$. Wir geben die Verknüpfungstabellen für $+$ und $\cdot$ auf $\mathbb{Z}/n\mathbb{Z}$ für $n = 4$ und $n = 7$ vollständig an.

Addition und Multiplikation auf $\mathbb{Z}/4\mathbb{Z}$

+	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{0}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{1}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{0}$
$\bar{2}$	$\bar{2}$	$\bar{3}$	$\bar{0}$	$\bar{1}$
$\bar{3}$	$\bar{3}$	$\bar{0}$	$\bar{1}$	$\bar{2}$

$\cdot$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$
$\bar{1}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{2}$	$\bar{0}$	$\bar{2}$	$\bar{0}$	$\bar{2}$
$\bar{3}$	$\bar{0}$	$\bar{3}$	$\bar{2}$	$\bar{1}$

Addition und Multiplikation auf $\mathbb{Z}/7\mathbb{Z}$

+	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$
$\bar{0}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$
$\bar{1}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{0}$
$\bar{2}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{0}$	$\bar{1}$
$\bar{3}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{0}$	$\bar{1}$	$\bar{2}$
$\bar{4}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$
$\bar{5}$	$\bar{5}$	$\bar{6}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$
$\bar{6}$	$\bar{6}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$

$\cdot$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$
$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$
$\bar{1}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$
$\bar{2}$	$\bar{0}$	$\bar{2}$	$\bar{4}$	$\bar{6}$	$\bar{1}$	$\bar{3}$	$\bar{5}$
$\bar{3}$	$\bar{0}$	$\bar{3}$	$\bar{6}$	$\bar{2}$	$\bar{5}$	$\bar{1}$	$\bar{4}$
$\bar{4}$	$\bar{0}$	$\bar{4}$	$\bar{1}$	$\bar{5}$	$\bar{2}$	$\bar{6}$	$\bar{3}$
$\bar{5}$	$\bar{0}$	$\bar{5}$	$\bar{3}$	$\bar{1}$	$\bar{6}$	$\bar{4}$	$\bar{2}$
$\bar{6}$	$\bar{0}$	$\bar{6}$	$\bar{5}$	$\bar{4}$	$\bar{3}$	$\bar{2}$	$\bar{1}$

Die Mengen $\mathbb{Z}/n\mathbb{Z}$ mit diesen beiden Verknüpfungen liefern uns neue Beispiele für Ringe, die nicht durch die bekannten Zahlbereiche gegeben sind und im Unterschied zu diesen aus nur endlich vielen Elementen bestehen.

(5.10) Satz Sei $n \in \mathbb{N}$. Dann bildet das Tripel $(\mathbb{Z}/n\mathbb{Z}, +, \cdot)$ einen Ring, mit $\bar{0} = 0 + n\mathbb{Z}$ als Null- und $\bar{1} = 1 + n\mathbb{Z}$ als Einselement. Man bezeichnet ihn als **Restklassenring** modulo n .

Beweis: Zunächst zeigen wir, dass $(\mathbb{Z}/n\mathbb{Z}, +)$ eine abelsche Gruppe ist, mit $\bar{0} = 0 + n\mathbb{Z}$ als Neutralelement. Zum Nachweis des Assoziativgesetzes seien $a + n\mathbb{Z}$, $b + n\mathbb{Z}$ und $c + n\mathbb{Z}$ vorgegeben, mit $a, b, c \in \mathbb{Z}$. Auf Grund der Definition der Verknüpfung $+$ auf $\mathbb{Z}/n\mathbb{Z}$, und auf Grund der Gültigkeit des Assoziativgesetzes für die Addition auf den ganzen Zahlen, erhalten wir

$$\begin{aligned} ((a + n\mathbb{Z}) + (b + n\mathbb{Z})) + (c + n\mathbb{Z}) &= ((a + b) + n\mathbb{Z}) + (c + n\mathbb{Z}) = ((a + b) + c) + n\mathbb{Z} = \\ (a + (b + c)) + n\mathbb{Z} &= (a + n\mathbb{Z}) + ((b + c) + n\mathbb{Z}) = (a + n\mathbb{Z}) + ((b + n\mathbb{Z}) + (c + n\mathbb{Z})) \end{aligned}$$

Für jedes Element $a + n\mathbb{Z}$ mit $a \in \mathbb{Z}$ gilt $(a + n\mathbb{Z}) + (0 + n\mathbb{Z}) = (a + 0) + n\mathbb{Z} = a + n\mathbb{Z}$ und $(0 + n\mathbb{Z}) + (a + n\mathbb{Z}) = (0 + a) + n\mathbb{Z} = a + n\mathbb{Z}$. Dies zeigt, dass $\bar{0} = 0_n \mathbb{Z}$ in der Halbgruppe $(\mathbb{Z}/n\mathbb{Z}, +)$ tatsächlich ein Neutralelement ist. Außerdem gilt jeweils $(a + n\mathbb{Z}) + ((-a) + n\mathbb{Z}) = (a + (-a)) + n\mathbb{Z} = 0 + n\mathbb{Z}$ und $((-a) + n\mathbb{Z}) + (a + n\mathbb{Z}) = ((-a) + a) + n\mathbb{Z} = 0 + n\mathbb{Z}$. Daran sehen wir, dass $(-a) + n\mathbb{Z}$ das Inverse von $a + n\mathbb{Z}$ im Monoid $(\mathbb{Z}/n\mathbb{Z}, +)$ ist. Schließlich ist auch das Kommutativgesetz erfüllt, denn für alle $a + n\mathbb{Z}, b + n\mathbb{Z} \in \mathbb{Z}/n\mathbb{Z}$ mit $a, b \in \mathbb{Z}$ gilt

$$(a + n\mathbb{Z}) + (b + n\mathbb{Z}) = (a + b) + n\mathbb{Z} = (b + a) + n\mathbb{Z} = (b + n\mathbb{Z}) + (a + n\mathbb{Z}).$$

Insgesamt ist $(\mathbb{Z}/n\mathbb{Z}, +)$ also tatsächlich eine abelsche Gruppe.

Als nächstes zeigen wir, dass durch $(\mathbb{Z}/n\mathbb{Z}, \cdot)$ ein abelsches Monoid gegeben ist. Zur Überprüfung von Kommutativ- und Assoziativgesetz seien $a + n\mathbb{Z}$, $b + n\mathbb{Z}$, $c + n\mathbb{Z}$ mit $a, b, c \in \mathbb{Z}$ vorgegeben. Dann gilt

$$(a + n\mathbb{Z}) \cdot (b + n\mathbb{Z}) = ab + n\mathbb{Z} = ba + n\mathbb{Z} = (b + n\mathbb{Z}) \cdot (a + n\mathbb{Z}).$$

und

$$\begin{aligned} ((a + n\mathbb{Z}) \cdot (b + n\mathbb{Z})) \cdot (c + n\mathbb{Z}) &= (ab + n\mathbb{Z}) \cdot (c + n\mathbb{Z}) = (ab)c + n\mathbb{Z} = \\ a(bc) + n\mathbb{Z} &= (a + n\mathbb{Z}) \cdot (bc + n\mathbb{Z}) = (a + n\mathbb{Z}) \cdot ((b + n\mathbb{Z}) \cdot (c + n\mathbb{Z})). \end{aligned}$$

Für jedes $a + n\mathbb{Z}$ mit $a \in \mathbb{Z}$ gilt außerdem $(a + n\mathbb{Z}) \cdot (1 + n\mathbb{Z}) = a \cdot 1 + n\mathbb{Z} = a + n\mathbb{Z}$ und $(1 + n\mathbb{Z}) \cdot (a + n\mathbb{Z}) = 1 \cdot a + n\mathbb{Z} = a + n\mathbb{Z}$. Dies zeigt, dass $\bar{1} = 1 + n\mathbb{Z}$ in der Halbgruppe $(\mathbb{Z}/n\mathbb{Z}, \cdot)$ ein Neutralelement ist. Insgesamt ist $(\mathbb{Z}/n\mathbb{Z}, \cdot)$ also ein abelsches Monoid.

Um zu zeigen, dass $(\mathbb{Z}/n\mathbb{Z}, +, \cdot)$ ein Ring ist, fehlt noch die Überprüfung des Distributivgesetzes. Seien dazu $a + n\mathbb{Z}$, $b + n\mathbb{Z}$, $c + n\mathbb{Z}$ vorgegeben, mit $a, b, c \in \mathbb{Z}$. Dann gilt

$$\begin{aligned} (a + n\mathbb{Z}) \cdot ((b + n\mathbb{Z}) + (c + n\mathbb{Z})) &= (a + n\mathbb{Z}) \cdot ((b + c) + n\mathbb{Z}) = a(b + c) + n\mathbb{Z} = \\ (ab + ac) + n\mathbb{Z} &= (ab + n\mathbb{Z}) + (ac + n\mathbb{Z}) = (a + n\mathbb{Z}) \cdot (b + n\mathbb{Z}) + (a + n\mathbb{Z}) \cdot (c + n\mathbb{Z}). \end{aligned}$$

Damit ist die Verifikation der Ringeigenschaften abgeschlossen. □

Aus den Ringaxiomen folgt unter anderem, dass jedes Element $a \in \mathbb{Z}/n\mathbb{Z}$ ein **Negatives** besitzt, also ein $b \in \mathbb{Z}/n\mathbb{Z}$ mit $a + b = \bar{0}$. Das Negative von $c + n\mathbb{Z}$ (mit $c \in \mathbb{Z}$) ist jeweils gegeben durch $(-c) + n\mathbb{Z} = (n - c) + n\mathbb{Z}$. In $\mathbb{Z}/7\mathbb{Z}$ gilt beispielsweise $-\bar{0} = \bar{0}$, $-\bar{1} = \bar{6}$, $-\bar{2} = \bar{5}$, $-\bar{3} = \bar{4}$, $-\bar{4} = \bar{3}$, $-\bar{5} = \bar{2}$ und $-\bar{6} = \bar{1}$. Dementsprechend lässt sich auf $\mathbb{Z}/n\mathbb{Z}$ eine **Subtraktion** definieren, indem man für $a, b \in \mathbb{Z}/n\mathbb{Z}$ jeweils $a - b = a + (-b)$ setzt. In $\mathbb{Z}/7\mathbb{Z}$ gilt beispielsweise $\bar{3} - \bar{4} = \bar{3} + (-\bar{4}) = \bar{3} + \bar{3} = \bar{6}$.

Eine naheliegende Frage lautet, ob $\mathbb{Z}/n\mathbb{Z}$ nicht vielleicht sogar ein Körper ist. Dazu müsste unter anderem gezeigt werden, dass jedes Element $\neq \bar{0}$ in $\mathbb{Z}/n\mathbb{Z}$ einen Kehrwert besitzt. Aber dies ist, zumindest für beliebiges n , nicht der Fall. In $\mathbb{Z}/4\mathbb{Z}$ gilt beispielsweise $\bar{2} \cdot \bar{0} = \bar{0}$, $\bar{2} \cdot \bar{1} = \bar{2}$, $\bar{2} \cdot \bar{2} = \bar{4} = \bar{0}$ und $\bar{2} \cdot \bar{3} = \bar{6} = \bar{2}$. Es existiert also kein $a \in \mathbb{Z}/4\mathbb{Z}$ mit $\bar{2} \cdot a = \bar{1}$, das Element $\bar{2}$ besitzt also keinen Kehrwert. Außerdem sehen wir, dass es in $\mathbb{Z}/4\mathbb{Z}$ Elemente ungleich $\bar{0}$ gibt, deren Produkt gleich $\bar{0}$ ist, nämlich $\bar{2} \cdot \bar{2} = \bar{0}$. Wie wir aus Teil (iv) von Proposition (5.8) wissen, ist dies in einem Körper nicht möglich.

Wir werden aber sehen, dass $\mathbb{Z}/n\mathbb{Z}$ in einige Fälle doch ein Körper ist. Um zu sehen, für welche natürlichen Zahlen n dies gilt, benötigen wir noch etwas Vorbereitung. Wir bezeichnen zwei natürliche Zahlen m und n als **teilerfremd**, wenn kein $d \in \mathbb{N}$ mit $d > 1$, $d \mid m$ und $d \mid n$ existiert.

(5.11) Lemma Sind $m, n \in \mathbb{N}$ teilerfremd, dann gibt es $a, b \in \mathbb{Z}$ mit $am + bn = 1$.

Beweis: Sei $S = \{am + bn \mid a, b \in \mathbb{Z}\}$. Zu zeigen ist, dass $1 \in S$ gilt. Sei $d \in \mathbb{N}$ die kleinste natürliche Zahl in S . Dann gilt $d \mid s$ für alle $s \in S$. Denn nehmen wir an, dies ist nicht der Fall, d.h. es gilt $d \nmid s$ für ein $s \in S$. Durch Division mit Rest erhalten wir dann $q, r \in \mathbb{Z}$ mit $s = qd + r$, wobei $0 < r < d$ gilt. Wir zeigen, dass r in S enthalten ist. Wegen $d \in S$ gibt es $a_0, b_0 \in \mathbb{Z}$ mit $d = a_0m + b_0n$. Wegen $s \in S$ existieren ebenso $a_1, b_1 \in \mathbb{Z}$ mit $s = a_1m + b_1n$. Es folgt

$$r = s - qd = (a_1m + b_1n) - q(a_0m + b_0n) = (a_1 - qa_0)m + (b_1 - qb_0)n,$$

wegen $a_1 - qa_0, b_1 - qb_0 \in \mathbb{Z}$ also tatsächlich $r \in S$. Aber $r \in S$, $r \in \mathbb{N}$ und $r < d$ stehen im Widerspruch zur Minimalität von d .

Damit ist gezeigt, dass $d \mid s$ für alle $s \in S$ erfüllt ist. Wegen $m, n \in S$ folgt $d \mid m$ und $d \mid n$. Weil m und n aber nach Voraussetzung teilerfremd sind, muss $d = 1$ sein. Damit ist $1 \in S$ nachgewiesen. Nach Definition von S existieren also $a, b \in \mathbb{Z}$ mit $1 = am + bn$. □

Wie in der Schulmathematik bezeichnen wir eine natürliche Zahl p als **Primzahl**, wenn $p > 1$ ist und keine $r, s \in \mathbb{N}$ mit $p = rs$ und $1 < r, s < p$ existieren, die Zahl also nicht in zwei echt kleinere Faktoren zerlegbar ist.

(5.12) Satz Für jedes $n \in \mathbb{N}$ ist der Restklassenring $\mathbb{Z}/n\mathbb{Z}$ genau dann ein Körper, wenn n eine Primzahl ist.

Beweis: „ $\Leftarrow$ “ Sei $p \in \mathbb{N}$ eine Primzahl. Wir zeigen, dass $\mathbb{Z}/p\mathbb{Z}$ ein Körper ist und müssen dazu nachweisen, dass jedes Element ungleich $\bar{0}$ in $\mathbb{Z}/p\mathbb{Z}$ einen Kehrwert besitzt. Sei $c + p\mathbb{Z}$ ein solches Element, mit $c \in \{1, \dots, p-1\}$. Weil p eine Primzahl ist, besitzt p in $\mathbb{N}$ nur die beiden Teiler 1 und p . Daraus folgt, dass c und p teilerfremd sind. Nach Lemma (5.11) existieren deshalb $a, b \in \mathbb{Z}$ mit $ac + bp = 1$. In $\mathbb{Z}/p\mathbb{Z}$ gilt deshalb die Gleichung

$$(a + p\mathbb{Z})(c + p\mathbb{Z}) + (b + p\mathbb{Z})(p + p\mathbb{Z}) = 1 + p\mathbb{Z}.$$

Wegen $p + p\mathbb{Z} = \bar{0}$ und $1 + p\mathbb{Z} = \bar{1}$ erhalten wir in $\mathbb{Z}/p\mathbb{Z}$ die Gleichung $(a + p\mathbb{Z})(c + p\mathbb{Z}) = \bar{1}$. Das Element $c + p\mathbb{Z}$ besitzt in $\mathbb{Z}/p\mathbb{Z}$ also das Element $a + p\mathbb{Z}$ als Kehrwert.

„ $\Rightarrow$ “ Ist n keine Primzahl, dann gilt entweder $n = 1$, oder es gibt $r, s \in \mathbb{N}$ mit $1 < r, s < n$ und $n = rs$. Wir zeigen, dass $\mathbb{Z}/n\mathbb{Z}$ in beiden Fällen kein Körper ist. Im Fall $n = 1$ gilt $\bar{1} = 1 + 1\mathbb{Z} = 0 + 1\mathbb{Z} = \bar{0}$, Null- und Einselement stimmen in $\mathbb{Z}/n\mathbb{Z}$ also überein. Wie wir oben gesehen haben, ist das in einem Körper ausgeschlossen. Betrachten wir nun noch den Fall $n = rs$ mit r und s wie oben angegeben. Setzen wir $a = r + n\mathbb{Z}$ und $b = s + n\mathbb{Z}$, dann gilt einerseits $a \neq \bar{0}$ und $b \neq \bar{0}$, andererseits aber $ab = (r + n\mathbb{Z})(s + n\mathbb{Z}) = rs + n\mathbb{Z} = n + n\mathbb{Z} = \bar{0}$. Nach Teil (iv) von Proposition (5.8) folgt daraus, dass $\mathbb{Z}/n\mathbb{Z}$ kein Körper ist. $\square$

Ist p eine Primzahl, dann verwendet man an Stelle von $\mathbb{Z}/p\mathbb{Z}$ für den Restklassenring auch die Notation $\mathbb{F}_p$. (Dabei steht das $\mathbb{F}$ für die englische Bezeichnung der algebraischen Struktur „Körper“. Diese lautet „field“.)

Beispielsweise ist $\mathbb{F}_7 = \mathbb{Z}/7\mathbb{Z}$ ein Körper, denn für jedes $a \in \mathbb{F}_7$ gibt es ein $b \in \mathbb{F}_7$ mit $ab = \bar{1}$: Es gilt $\bar{1} \cdot \bar{1} = \bar{1}$, $\bar{6} \cdot \bar{6} = \bar{1}$, $\bar{2} \cdot \bar{4} = \bar{1}$ und $\bar{3} \cdot \bar{5} = \bar{1}$. Insgesamt sind die Kehrwerte in $\mathbb{F}_7$ also durch folgende Tabelle gegeben:

a	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$
a^{-1}	—	$\bar{1}$	$\bar{4}$	$\bar{5}$	$\bar{2}$	$\bar{3}$	$\bar{6}$

Für den Körper $\mathbb{F}_{13}$ erhält man auf gleiche Weise

a	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{7}$	$\bar{8}$	$\bar{9}$	$\bar{10}$	$\bar{11}$	$\bar{12}$
a^{-1}	—	$\bar{1}$	$\bar{7}$	$\bar{9}$	$\bar{10}$	$\bar{8}$	$\bar{11}$	$\bar{2}$	$\bar{5}$	$\bar{3}$	$\bar{4}$	$\bar{6}$	$\bar{12}$

Bisher haben wir nur Strukturen kennengelernt, in denen die zugehörigen Verknüpfungen kommutativ sind. Bei den Objekten, die wir nun einführen, ist das Kommutativgesetz in der Regel nicht erfüllt. Sie werden später für uns ein wichtiges Hilfsmittel bei der Untersuchung und Lösung von Linearen Gleichungssystemen sein.

(5.13) Definition Seien $m, n \in \mathbb{N}$, und sei R ein Ring. Eine **$m \times n$ - Matrix** über R ist eine Abbildung

$$A: \{1, \dots, m\} \times \{1, \dots, n\} \longrightarrow R.$$

Dabei nennt man $A(i, j)$ den **Eintrag** von A an der Stelle (i, j) . Die Menge aller $m \times n$ -Matrizen über R wird mit $\mathcal{M}_{m \times n, R}$ bezeichnet. An Stelle von $\mathcal{M}_{m \times n, R}$ schreibt man auch kürzer $\mathcal{M}_{n, R}$. Die Elemente dieser Menge werden als **quadratische** Matrizen bezeichnet.

Durch die Gleichung $A = (a_{ij})$ ordnet man dem Eintrag $A(i, j)$ der Matrix A die Bezeichnung a_{ij} zu. Allgemein legen wir folgende Konvention fest: Bezeichnet ein Großbuchstabe wie zum Beispiel A, B, C eine Matrix, dann bezeichnen die indizierten Kleinbuchstaben a_{ij} , b_{ij} , c_{ij} immer automatisch die Einträge dieser Matrix. Man kann eine Matrix auch auf übersichtliche Weise als rechteckiges Schema der Form

$$A = \begin{pmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ & & \dots & \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{pmatrix} \quad \text{darstellen.}$$

Allgemein nennt man $a_{i\bullet} = (a_{i1}, \dots, a_{in}) \in R^n$ die ***i-te Zeile*** und $a_{\bullet j} = (a_{1j}, \dots, a_{mj}) \in R^m$ die ***j-te Spalte*** von A .

Offenbar existiert eine natürliche Bijektion zwischen Matrizen mit nur einer Zeile (also Elementen aus $\mathcal{M}_{1 \times n, R}$), Matrizen mit nur einer Spalte (Elementen der Menge $\mathcal{M}_{n \times 1, R}$) und dem R^n . Dabei entspricht $(a_1, \dots, a_n) \in R^n$ den beiden Matrizen

$$\begin{pmatrix} a_1 & a_2 & \dots & a_n \end{pmatrix} \in \mathcal{M}_{1 \times n, R} \quad \text{und} \quad \begin{pmatrix} a_1 \\ \vdots \\ a_n \end{pmatrix} \in \mathcal{M}_{n \times 1, R}.$$

Zur Beschreibung der Einträge verwendet man häufig als hilfreiche Abkürzung das sogenannte **Kronecker-Delta**. Für jeden Ring R und beliebige $m, n \in \mathbb{N}$ definiert man

$$\delta_{mn} = \delta_{mn, R} = \begin{cases} 1_R & \text{falls } m = n \\ 0_R & \text{falls } m \neq n. \end{cases}$$

Falls aus dem Kontext geschlossen werden kann, welcher Ring gemeint ist, wird der Index R auch oft weggelassen. Die folgenden konkreten Beispiele für Matrizen werden uns in den Anwendungen immer wieder begegnen.

- (i) die **Nullmatrix** $0^{(m \times n)}$ in $\mathcal{M}_{m \times n, R}$, deren sämtliche Einträge gleich 0_R sind
(Mit $0^{(n)} = 0^{(n \times n)}$ bezeichnen wir die quadratische Nullmatrix.)
- (ii) die **Einheitsmatrix** $E = E^{(n)}$ in $\mathcal{M}_{n, R}$ mit den Einträgen δ_{ij} für $1 \leq i \leq m$ und $1 \leq j \leq n$
(Die Einheitsmatrix ist also immer quadratisch.)
- (iii) die **Basismatrizen** $B_{k\ell} = B_{k\ell}^{(m \times n)}$ mit den Einträgen $b_{ij} = \delta_{ik} \delta_{j\ell}$ für $1 \leq i \leq m$ und $1 \leq j \leq n$

Beispielsweise ist

$$0^{(2 \times 3)} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}, \quad E^{(3)} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \text{und} \quad B_{12}^{(3 \times 2)} = \begin{pmatrix} 0 & 1 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}$$

Wir definieren nun eine Reihe von Rechenoperationen auf Matrizen. Weiterhin sei R ein beliebiger Ring. Um die Rechenoperationen definieren zu können, verwenden wir zum ersten Mal das **Summenzeichen** als Hilfsmittel. In allgemeiner Form werden wir dieses erst in § 11 einführen. Momentan genügt es zu wissen, dass für beliebige Ringelemente $a_1, \dots, a_n \in R$ der Ausdruck

$$\sum_{k=1}^n a_k \quad \text{eine Kurzschreibweise für} \quad a_1 + \dots + a_n \quad \text{ist.}$$

Den Buchstaben k bezeichnet man dabei als *Laufindex* der Summe. Welchen Buchstaben man dafür verwendet, spielt letztlich keine Rolle; beispielsweise haben $\sum_{j=1}^n a_j$ und $\sum_{k=1}^n a_k$ denselben Wert. Ist $A = (a_{ij})$ eine Matrix in $\mathcal{M}_{m \times n, R}$, so kann über eine Zeile oder eine Spalte summiert werden: Für $1 \leq k \leq m$ und $1 \leq \ell \leq n$ gilt

$$\sum_{j=1}^n a_{kj} = a_{k1} + \dots + a_{kn} \quad \text{und} \quad \sum_{i=1}^m a_{i\ell} = a_{1\ell} + \dots + a_{m\ell}.$$

Der erste Ausdruck ist die Summe über die k -te Zeile von A , und beim zweiten Ausdruck handelt es sich um die Summe über die ℓ -te Spalte. Kommen wir nun zur Definition der Rechenoperationen für Matrizen.

- (i) Seien $m, n \in \mathbb{N}$ und $A, B \in \mathcal{M}_{m \times n, R}$ mit Einträgen $A = (a_{ij})$ und $B = (b_{ij})$. Dann nennt man die Matrix $C = (c_{ij})$ mit $c_{ij} = a_{ij} + b_{ij}$ für $1 \leq i \leq m$ und $1 \leq j \leq n$ die **Summe** von A und B . Wir bezeichnen diese Matrix mit $A + B$. Beispielsweise gilt

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} + \begin{pmatrix} -1 & 5 & -2 \\ 3 & 0 & 1 \end{pmatrix} = \begin{pmatrix} 0 & 7 & 1 \\ 7 & 5 & 7 \end{pmatrix}.$$

- (ii) Sei $A \in \mathcal{M}_{m \times n, R}$ mit $A = (a_{ij})$ und $\lambda \in K$. Dann ist die Matrix $C = (c_{ij})$ mit $c_{ij} = \lambda a_{ij}$ ein **skalares Vielfaches** von A , das wir mit λA bezeichnen. Beispielsweise ist

$$7 \cdot \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} = \begin{pmatrix} 7 & 14 & 21 \\ 28 & 35 & 42 \end{pmatrix}.$$

- (iii) Seien nun $m, n, r \in \mathbb{N}$ und $A \in \mathcal{M}_{m \times n, R}$, $B \in \mathcal{M}_{n \times r, R}$. Dann heißt die Matrix $C \in \mathcal{M}_{m \times r, R}$ mit den Einträgen

$$c_{ij} = \sum_{k=1}^n a_{ik} b_{kj}$$

Produkt der Matrizen A und B und wird mit AB bezeichnet. Auch hierzu ein konkretes Beispiel:

$$\begin{pmatrix} 1 & 0 \\ -1 & 1 \\ 3 & 0 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 & 4 \\ 4 & 3 & 2 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 3 & 1 & -1 & -3 \\ 3 & 6 & 9 & 12 \end{pmatrix}$$

Den Eintrag an der Position $(2, 2)$ erhält man zum Beispiel durch Multiplikation der zweiten Zeile der ersten Matrix mit der zweiten Spalte der zweiten Matrix, also durch die Rechnung $(-1) \cdot 2 + 1 \cdot 3 = 1$. Der Eintrag an der Position $(3, 3)$ kommt entsprechend durch $3 \cdot 3 + 0 \cdot 2 = 9$ zu Stande.

- (iv) Sei $A \in \mathcal{M}_{m \times n, R}$. Die Matrix $B \in \mathcal{M}_{n \times m, R}$ mit den Einträgen $b_{ij} = a_{ji}$ für $1 \leq i \leq n$ und $1 \leq j \leq m$ wird die zu A **transponierte** Matrix tA genannt. Zum Beispiel ist

$${}^t \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{pmatrix} = \begin{pmatrix} 1 & 4 \\ 2 & 5 \\ 3 & 6 \end{pmatrix}.$$

Um den Eintrag c_{ij} der Produktmatrix $C = AB$ an der Position (i, j) zu erhalten, muss die i -te Zeile der Matrix A mit der j -ten Spalte der Matrix B multipliziert werden. Man beachte, dass das Produkt AB nur gebildet werden kann, wenn die Spaltenzahl von A mit der Zeilenzahl von B übereinstimmt. Die Summe $A + B$ ist nur dann definiert, wenn A und B dasselbe Format, also dieselbe Anzahl Zeilen und Spalten besitzen.

Die neu eingeführten Rechenoperationen erfüllen eine Reihe von Rechenregeln. Wir beginnen mit der Summe von Matrizen.

(5.14) Proposition Sei R ein Ring, und seien $m, n \in \mathbb{N}$. Dann bildet die Menge $\mathcal{M}_{m \times n, R}$ mit der Addition von Matrizen eine abelsche Gruppe. Dabei ist $0^{(m \times n)}$ das Neutralelement, und für jedes $A \in \mathcal{M}_{m \times n, R}$ ist $(-1_R)A$ das Inverse von A . Dieses Inverse wird das **Negative** von A genannt und mit $-A$ bezeichnet.

Beweis: Zunächst zeigen wir, dass die Addition von Matrizen das Assoziativ- und das Kommutativgesetz erfüllt. Seien dazu $A, B, C \in \mathcal{M}_{m \times n, R}$ vorgegeben. Für $1 \leq i \leq m$ und $1 \leq j \leq n$ stimmt wegen $(a_{ij} + b_{ij}) + c_{ij} = a_{ij} + (b_{ij} + c_{ij})$ der Eintrag der Matrix $(A + B) + C$ an der Stelle (i, j) überein mit dem Eintrag der Matrix $A + (B + C)$ an derselben Position.

Ebenso zeigt die Gleichung $a_{ij} + b_{ij} = b_{ij} + a_{ij}$, dass die Matrizen $A + B$ und $B + A$ an der Position (i, j) übereinstimmen. Insgesamt sind damit die Gleichungen $(A + B) + C = A + (B + C)$ und $A + B = B + A$ bewiesen. Das Paar $(\mathcal{M}_{m \times n, R}, +)$ ist also eine abelsche Halbgruppe.

Die Matrizen $A + 0^{(m \times n)}$ und $0^{(m \times n)} + A$ haben an der Position (i, j) jeweils den Eintrag $a_{ij} + 0_R = 0_R + a_{ij} = a_{ij}$. Es gilt also $A + 0^{(m \times n)} = 0^{(m \times n)} + A = A$. Dies zeigt, dass die Nullmatrix in der Halbgruppe $(\mathcal{M}_{m \times n, R}, +)$ ein Neutralelement ist und somit ein abelsches Monoid vorliegt. Die Matrix $(-1_R)A$ hat an der Stelle (i, j) den Eintrag $(-1_R)a_{ij} = -a_{ij}$. Dadurch ist die Bezeichnung $-A$ für diese Matrix gerechtfertigt. Die Matrizen $A + (-A)$ und $(-A) + A$ haben an der Position (i, j) beide den Eintrag $a_{ij} + (-a_{ij}) = (-a_{ij}) + a_{ij} = 0_R$. Dies zeigt, dass $-A$ im Monoid $(\mathcal{M}_{m \times n, R}, +)$ ein zu A inverses Element ist. Jedes Element in diesem Monoid ist also invertierbar. Dies zeigt insgesamt, dass $(\mathcal{M}_{m \times n, R}, +)$ eine abelsche Gruppe ist. $\square$

(5.15) Proposition Sei R ein Ring, und seien $m, n, r, s \in \mathbb{N}$. Weiter seien $A, A' \in \mathcal{M}_{m \times n, R}$, $B, B' \in \mathcal{M}_{n \times r, R}$ und $C \in \mathcal{M}_{r \times s, R}$. Dann gelten die folgenden Rechenregeln.

$$\begin{aligned} \text{(i)} \quad & A(B + B') = AB + AB' \text{ und } (A + A')B = AB + A'B & \text{(ii)} \quad & A(\lambda B) = (\lambda A)B = \lambda(AB) \\ \text{(iii)} \quad & E^{(m)}A = AE^{(n)} = A & \text{(iv)} \quad & (AB)C = A(BC) & \text{(v)} \quad & {}^t(AB) = {}^tB {}^tA \end{aligned}$$

Beweis: zu (i) Wir beschränken uns auf den Beweis der zweiten Gleichung, da der Beweis der ersten vollkommen analog verläuft. Es sei $C' = A + A'$, $D = C'B$, $F = AB$, $G = A'B$ und $H = F + G$. Dann ist $D = H$ zu zeigen. Für $1 \leq i \leq m$ und $1 \leq j \leq n$ gilt jeweils $c'_{ij} = a_{ij} + a'_{ij}$. Die Einträge d_{ij} der Matrix D (für $1 \leq i \leq m$, $1 \leq j \leq r$) sind gegeben durch

$$d_{ij} = \sum_{k=1}^n c'_{ik} b_{kj} = \sum_{k=1}^n (a_{ik} + a'_{ik}) b_{kj} = \sum_{k=1}^n a_{ik} b_{kj} + \sum_{k=1}^n a'_{ik} b_{kj}.$$

Für die Einträge f_{ij} und g_{ij} der Matrizen F und G (mit $1 \leq i \leq m$, $1 \leq j \leq r$) erhalten wir

$$f_{ij} = \sum_{k=1}^n a_{ik} b_{kj} \quad \text{und} \quad g_{ij} = \sum_{k=1}^n a'_{ik} b_{kj}.$$

Es folgt

$$h_{ij} = f_{ij} + g_{ij} = \sum_{k=1}^n a_{ik} b_{kj} + \sum_{k=1}^n a'_{ik} b_{kj} = d_{ij}$$

für $1 \leq i \leq m$, $1 \leq j \leq r$, also insgesamt $D = H$.

zu (ii) Wir definieren $C' = \lambda B$, $D = AC'$, $F = \lambda A$, $G = FB$, $H = AB$ und $U = \lambda H$. Zu zeigen ist dann $D = G = U$. Nach Definition gilt $c'_{ij} = \lambda b_{ij}$ für $1 \leq i \leq n$, $1 \leq j \leq r$ und

$$d_{ij} = \sum_{k=1}^n a_{ik} c'_{kj} = \sum_{k=1}^n \lambda a_{ik} b_{kj}$$

für $1 \leq i \leq m$ und $1 \leq j \leq r$. Andererseits gilt auch $f_{ij} = \lambda a_{ij}$ für $1 \leq i \leq m, 1 \leq j \leq n$ und $g_{ij} = \sum_{k=1}^n f_{ik} b_{kj} = \sum_{k=1}^n \lambda a_{ik} b_{kj} = d_{ij}$ für $1 \leq i \leq m, 1 \leq j \leq r$, womit die Gleichung $D = G$ bewiesen ist. Nun gilt außerdem $h_{ij} = \sum_{k=1}^n a_{ik} b_{kj}$ für $1 \leq i \leq m$ und $1 \leq j \leq r$, und

$$u_{ij} = \lambda h_{ij} = \sum_{k=1}^n \lambda a_{ik} b_{kj} = d_{ij}$$

für dieselben Paare (i, j) , wodurch auch die Gleichung $U = D$ bewiesen ist.

zu (iii) Wir beschränken uns auf den Beweis der Gleichung $E^{(m)}A = A$. Bezeichnen wir das Matrixprodukt $E^{(m)}A \in \mathcal{M}_{m \times n, K}$ mit B , dann ist der Eintrag $b_{k\ell}$ von B an der Position (k, ℓ) für $1 \leq k \leq m$ und $1 \leq \ell \leq n$ gegeben durch

$$b_{k\ell} = \sum_{j=1}^m \delta_{kj} a_{j\ell} = \delta_{kk} a_{k\ell} = a_{k\ell}.$$

Dies zeigt, dass B mit der Matrix A übereinstimmt.

zu (iv) Wir definieren $D = AB$, $F = DC$, $G = BC$ und $H = AG$. Dann gilt für $1 \leq k \leq m$ und $1 \leq \ell \leq r$ jeweils $d_{k\ell} = \sum_{i=1}^n a_{ki} b_{i\ell}$, und für die Einträge der Matrix F erhalten wir

$$f_{k\ell} = \sum_{i=1}^r d_{ki} c_{i\ell} = \sum_{i=1}^r \sum_{j=1}^n a_{kj} b_{ji} c_{i\ell},$$

für $1 \leq k \leq m$ und $1 \leq \ell \leq s$. Andererseits hat G die Einträge $g_{k\ell} = \sum_{i=1}^r b_{ki} c_{i\ell}$ für $1 \leq k \leq n$ und $1 \leq \ell \leq s$, und für die Einträge $h_{k\ell}$ der Matrix H ($1 \leq k \leq m, 1 \leq \ell \leq s$) gilt

$$h_{k\ell} = \sum_{i=1}^n a_{ki} g_{i\ell} = \sum_{i=1}^n \sum_{j=1}^r a_{ki} b_{ij} c_{j\ell}, \quad \text{also insgesamt } F = H.$$

zu (v) Hier definieren wir die Hilfsmatrizen $C = AB$, $D = {}^t C$, $F = {}^t A$, $G = {}^t B$ und $H = GF$. Dann müssen wir $D = H$ nachrechnen. Es gilt Für $1 \leq i \leq m$ und $1 \leq j \leq r$ gilt

$$c_{ij} = \sum_{k=1}^n a_{ik} b_{kj} \text{ für } 1 \leq i \leq m, 1 \leq j \leq r \quad \text{und} \quad d_{ij} = c_{ji} = \sum_{k=1}^n a_{jk} b_{ki} \text{ für } 1 \leq i \leq r, 1 \leq j \leq m.$$

Wegen $f_{ij} = a_{ji}$ und $g_{ij} = b_{ji}$ für $1 \leq i \leq n$ und $1 \leq j \leq m$ bzw. $1 \leq i \leq r$ und $1 \leq j \leq n$ gilt außerdem

$$h_{ij} = \sum_{k=1}^n g_{ik} f_{kj} = \sum_{k=1}^n b_{ki} a_{jk} = d_{ij}$$

für $1 \leq i \leq r$ und $1 \leq j \leq m$, also $H = D$ wie gewünscht. $\square$

(5.16) Folgerung Sei R ein Ring, und seien $n \in \mathbb{N}$. Dann bildet die Menge $\mathcal{M}_{n,R}$ mit der Multiplikation von Matrizen ein Monoid, mit $E^{(n)}$ als Neutralelement.

Beweis: Nach Teil (iv) von Proposition (5.15) ist die Multiplikation von Matrizen eine assoziative Verknüpfung auf $\mathcal{M}_{n,R}$. Aus Teil (iii) folgt, dass $E^{(n)}$ ein Neutralelement in der Halbgruppe $(\mathcal{M}_{n,R}, \cdot)$ ist. Insgesamt handelt es sich bei $(\mathcal{M}_{n,R}, \cdot)$ also um ein Monoid. $\square$

Das Monoid $(\mathcal{M}_{n,R}, \cdot)$ ist in der Regel nicht kommutativ. Ist beispielsweise $n = 2$, und stimmen im Ring R das Nullelement 0_R und das Einselement 1_R nicht überein, dann zeigt das Beispiel

$$\begin{pmatrix} 0_R & 1_R \\ 1_R & 0_R \end{pmatrix} \begin{pmatrix} 1_R & 1_R \\ 0_R & 1_R \end{pmatrix} = \begin{pmatrix} 0_R & 1_R \\ 1_R & 1_R \end{pmatrix} \quad , \quad \begin{pmatrix} 1_R & 1_R \\ 0_R & 1_R \end{pmatrix} \begin{pmatrix} 0_R & 1_R \\ 1_R & 0_R \end{pmatrix} = \begin{pmatrix} 1_R & 1_R \\ 1_R & 0_R \end{pmatrix}$$

dass die Gleichung $AB = BA$ nicht für alle $A, B \in \mathcal{M}_{n,R}$ erfüllt ist.

Wir bemerken auch, dass das Tupel $(\mathcal{M}_{n,R}, +, \cdot)$ alle Eigenschaften eines Rings besitzt, mit Ausnahme des Kommutativgesetzes der Multiplikation. Eine solche Struktur wird in der Algebra als **Schieferring** bezeichnet.

Eine Matrix $A \in \mathcal{M}_{n,R}$ wird **invertierbar** genannt, wenn eine Matrix $B \in \mathcal{M}_{n,R}$ mit $AB = BA = E^{(n)}$ existiert. Wir werden in einem späteren Kapitel sehen, wie man herausfindet, ob eine Matrix invertierbar ist, und wie sich gegebenenfalls das Inverse berechnen lässt.

(5.17) Folgerung Die Menge der invertierbaren $n \times n$ -Matrizen über einem Körper K bildet mit der Multiplikation von Matrizen eine Gruppe. Man nennt sie die **allgemeine lineare Gruppe** und bezeichnet sie mit $GL_n(K)$.

Beweis: Dies erhält man unmittelbar durch Anwendung von Folgerung (5.6) auf das Monoid $(\mathcal{M}_{n,R}, \cdot)$. □

§ 6. Vektorräume, lineare Abbildungen und lineare Gleichungssysteme

Inhaltsübersicht

Der Begriff des Vektorraums über einem Körper K ist für die Lineare Algebra von zentraler Bedeutung. Es handelt sich dabei um eine Menge V mit einer Verknüpfung, der *Vektoraddition*, und einer Abbildung $K \times V \rightarrow V$, der *skalaren Multiplikation*. Das wichtigste Beispiel sind die Vektorräume der Form K^n , mit $n \in \mathbb{N}$. Aber auch viele weitere Objekte der Mathematik besitzen eine Vektorraumstruktur, zum Beispiel die Menge $\mathcal{M}_{m \times n, K}$ der $(m \times n)$ -Matrizen über einem Körper K . Auf gewissen Teilmengen eines Vektorraums, den sog. *Untervektorräumen*, lässt sich ebenfalls eine Vektorraumstruktur definieren.

Eine Abbildung $V \rightarrow W$ zwischen Vektorräumen bezeichnet man als *lineare Abbildung*, wenn sie „verträglich“ mit der Vektoraddition und der skalaren Multiplikation der beiden Vektorräume ist. Wichtigstes Beispiel einer linearen Abbildung ist für uns zunächst die Matrix-Vektor-Multiplikation. Oft haben lineare Abbildungen eine geometrische Interpretation; zum Beispiel ist die Spiegelung im $\mathbb{R}^2$ an einer Geraden durch den Koordinatenursprung $(0,0)$ eine lineare Abbildung, ebenso jede Drehung um den Ursprung.

Desweiteren führen wir in diesem Kapitel den Begriff des *linearen Gleichungssystems* ein. Wir werden sehen, wie sich solche Systeme und ihre Lösungsmengen mit Hilfe der Untervektorräume und der linearen Abbildungen auf übersichtliche Weise beschreiben lassen. Dies dient auch zur Vorbereitung des nächsten Kapitels, wo wir uns dann mit der Berechnung von Lösungsmengen linearer Gleichungssysteme befassen.

Wichtige Begriffe und Sätze

- Vektorraum über einem Körper K
- Untervektorraum, affiner Unterraum
- lineare Abbildung, affin-lineare Abbildung, Affinität
- allgemeine lineare Gruppe $GL(V)$ zu einem Vektorraum V
- Kern und Bild einer linearen Abbildung
- homogenes und inhomogenes lineares Gleichungssystem
- (erweiterte) Koeffizientenmatrix eines linearen Gleichungssystems
- Lösungsmenge eines linearen Gleichungssystems
- (eindeutige) Lösbarkeit eines linearen Gleichungssystems
- Vektorraum $\text{Hom}_K(V, W)$ der linearen Abbildungen $V \rightarrow W$

(6.1) Definition Sei K ein Körper. Ein K -**Vektorraum** ist ein Tripel $(V, +, \cdot)$ bestehend aus einer nichtleeren Menge V und Abbildungen $+: V \times V \rightarrow V$ und $\cdot: K \times V \rightarrow V$ genannt **Vektoraddition** und **skalare Multiplikation**, so dass folgende Bedingungen erfüllt sind.

- (i) Das Paar $(V, +)$ ist eine abelsche Gruppe.
- (ii) Für alle $v, w \in V$ und $\lambda, \mu \in K$ gelten die Rechenregeln

$$(a) \quad (\lambda + \mu) \cdot v = (\lambda \cdot v) + (\mu \cdot v)$$

$$(b) \quad \lambda \cdot (v + w) = (\lambda \cdot v) + (\lambda \cdot w)$$

$$(c) \quad (\lambda \mu) \cdot v = \lambda \cdot (\mu \cdot v)$$

$$(d) \quad 1_K \cdot v = v$$

Die Elemente der Menge V werden **Vektoren** genannt.

Bei der skalaren Multiplikation wird häufig auf das Abbildungssymbol $\cdot$ verzichtet. Das Neutralelement der Gruppe $(V, +)$ bezeichnet man als den **Nullvektor** 0_V des Vektorraums. Das Inverse eines Vektors $v \in V$ bezüglich der Vektoraddition bezeichnet man mit $-v$ und verwendet $v - w$ als abkürzende Schreibweise für $v + (-w)$. Per Konvention bindet die skalare Multiplikation stärker als die Vektoraddition, d.h. der Ausdruck $\lambda v + w$ ist gleichbedeutend mit $(\lambda v) + w$ für $\lambda \in K, v, w \in V$.

(6.2) Proposition Sei K ein Körper. Die folgenden Strukturen sind Beispiele für K -Vektorräume.

- (i) das Tripel $(K^n, +, \cdot)$ ($n \in \mathbb{N}$), wobei die Abbildung $+: K^n \times K^n \rightarrow K^n$ und $\cdot: K \times K^n \rightarrow K^n$ definiert sind durch

$$(a_1, \dots, a_n) + (b_1, \dots, b_n) = (a_1 + b_1, \dots, a_n + b_n) \quad \text{und} \quad \lambda \cdot (a_1, \dots, a_n) = (\lambda a_1, \dots, \lambda a_n).$$

Insbesondere ist $K^1 = K$ selbst ein K -Vektorraum. Hier ist die Vektoraddition die Addition auf K , und die skalare Multiplikation ist die Multiplikation auf K .

- (ii) das Tripel $(\{0_K\}, +, \cdot)$, mit den Abbildungen $+: \{0_K\} \times \{0_K\} \rightarrow \{0_K\}$ und $\cdot: K \times \{0_K\} \rightarrow \{0_K\}$ gegeben durch $0_K + 0_K = 0_K$ und $\lambda \cdot 0_K = 0_K$ für alle $\lambda \in K$
- (iii) das Tripel $(\mathcal{M}_{m \times n, K}, +, \cdot)$ ($m, n \in \mathbb{N}$) wobei $+$ die Addition von Matrizen und $\cdot: K \times \mathcal{M}_{m \times n, K} \rightarrow \mathcal{M}_{m \times n, K}$ durch $(\lambda, A) \mapsto \lambda A$ gegeben ist
- (iv) Jeder $\mathbb{C}$ -Vektorraum $(V, +, \cdot)$ besitzt auch eine Struktur als $\mathbb{R}$ -Vektorraum, gegeben durch $(V, +, \cdot_{\mathbb{R}})$, wobei die Abbildung $\cdot_{\mathbb{R}}: \mathbb{R} \times V \rightarrow V$ durch Einschränkung der Abbildung $\cdot: \mathbb{C} \times V \rightarrow V$ auf $\mathbb{R} \times V$ zu Stande kommt.

Beweis: In jedem Fall können die Vektorraum-Eigenschaften unmittelbar überprüft werden. Wir beschränken uns auf den Nachweis im Fall (iii). Aus Proposition (5.14) wissen wir bereits, dass $(\mathcal{M}_{m \times n, K}, +)$ eine abelsche Gruppe ist.

Zum Nachweis der übrigen Regeln seien $A, B \in \mathcal{M}_{m \times n, K}$ und $\lambda, \mu \in K$ vorgegeben. Es gilt $(\lambda + \mu)A = \lambda A + \mu A$, denn für $1 \leq i \leq m$ und $1 \leq j \leq n$ ist der Eintrag der Matrix an der Position (i, j) jeweils gleich $(\lambda + \mu)a_{ij} = \lambda a_{ij} + \mu a_{ij}$. Es gilt $\lambda(A + B) = \lambda A + \lambda B$, weil der Eintrag an der Stelle j jeweils mit $\lambda(a_{ij} + b_{ij}) = \lambda a_{ij} + \lambda b_{ij}$ übereinstimmt. Beide Matrixgleichungen ergeben sich also direkt aus dem Distributivgesetz des Körpers K . Die Gleichung $(\lambda\mu)A = \lambda(\mu A)$ ergibt sich direkt aus dem Assoziativgesetz der multiplikativen Verknüpfung des Körpers K , denn für den Eintrag an der Stelle (i, j) gilt jeweils $(\lambda\mu)a_{ij} = \lambda(\mu a_{ij})$. Schließlich gilt auch $1_K \cdot A = A$, denn der Eintrag der Matrix an der Stelle (i, j) ist gleich $1_K \cdot a_{ij} = a_{ij}$. $\square$

(6.3) Lemma Sei $(V, +, \cdot)$ ein K -Vektorraum. Dann gilt für alle $\lambda \in K$ und $v \in V$ die Äquivalenz

$$\lambda v = 0_V \iff \lambda = 0_K \text{ oder } v = 0_V ,$$

außerdem $(-1_K)v = -v$ für alle $v \in V$.

Beweis: Zunächst beweisen wir die Äquivalenz. „ $\Leftarrow$ “ Ist $\lambda = 0_K$, dann gilt $\lambda v = 0_K v = (0_K + 0_K)v = 0_K v + 0_K v = \lambda v + \lambda v$. Addition von $-\lambda v$ auf beiden Seiten dieser Gleichung liefert

$$\lambda v + (-\lambda v) = \lambda v + \lambda v + (-\lambda v) \iff 0_V = \lambda v + 0_V \iff 0_V = \lambda v.$$

Setzen wir nun $v = 0_V$ voraus, dann erhalten wir $\lambda v = \lambda 0_V = \lambda(0_V + 0_V) = \lambda 0_V + \lambda 0_V = \lambda v + \lambda v$. Wieder führt die Addition von $-\lambda v$ auf beiden Seiten zum gewünschten Ergebnis.

„ $\Rightarrow$ “ Setzen wir $\lambda v = 0_V$ voraus, und nehmen wir an, es ist $\lambda \neq 0_K$. Dann gilt

$$v = 1_K v = (\lambda^{-1} \lambda) v = \lambda^{-1}(\lambda v) = \lambda^{-1} 0_V = 0_V ,$$

wobei im letzten Schritt die bereits bewiesene Rechenregel $\mu 0_V = 0_V$ für alle $\mu \in K$ verwendet wurde. Beweisen wir nun noch die Gleichung $(-1_K)v = -v$. Es gilt $v + (-1_K)v = 1_K v + (-1_K)v = (1_K + (-1_K))v = 0_K v = 0_V$. Addition von $-v$ auf beiden Seiten liefert

$$v + (-1_K)v + (-v) = 0_V + (-v) \iff v + (-v) + (-1_K)v = -v \iff$$

$$0_V + (-1_K)v = -v \iff (-1_K)v = -v$$

$\square$

(6.4) Definition Sei V ein K -Vektorraum. Eine Teilmenge $U \subseteq V$ wird **Untervektorraum** von V genannt, wenn folgende Bedingungen erfüllt sind.

- (i) $0_V \in U$
- (ii) $v + w \in U$ für alle $v, w \in U$
- (iii) $\lambda v \in U$ für alle $\lambda \in K$ und $v \in U$

Motiviert ist die Definition des Untervektorraumbegriffs durch den Wunsch, auf gewissen Teilmengen eines Vektorraums ebenfalls eine Vektorraumstruktur zu erhalten. Der folgende Satz zeigt, dass die Definition die gewünschte Funktion erfüllt.

(6.5) Satz Sei $(V, +, \cdot)$ ein K -Vektorraum und $U \subseteq V$ ein Untervektorraum von V . Definieren wir Abbildungen $+_U : U \times U \rightarrow V$ und $\cdot_U : K \times U \rightarrow V$ durch

$$v +_U w = v + w \quad \text{und} \quad \lambda \cdot_U v = \lambda \cdot v \quad \text{für } v, w \in U \text{ und } \lambda \in K,$$

dann ist durch $(U, +_U, \cdot_U)$ ein K -Vektorraum gegeben.

Beweis: Weil U ein Untervektorraum ist, gilt $v + w \in U$ für alle $v, w \in U$, also auch $v +_U w = v + w \in U$. Dies zeigt, dass $+_U$ eine Abbildung $U \times U \rightarrow U$, also eine Verknüpfung auf U gegeben ist. Ebenso gilt $\lambda \cdot_U v = \lambda \cdot v \in U$ für alle $v \in U$ und $\lambda \in K$. Somit ist $\cdot_U$ eine Abbildung $K \times U \rightarrow U$. Wir müssen nun überprüfen, dass $(U, +_U, \cdot_U)$ die Vektorraum-Bedingungen aus Definition (6.1) erfüllt.

Zunächst überprüfen wir, dass $(U, +_U)$ eine abelsche Gruppe ist. Für alle $u, v, w \in U$ gilt $(u +_U v) +_U w = (u + v) + w = u + (v + w) = u +_U (v +_U w)$, also ist das Assoziativgesetz erfüllt. Nach Voraussetzung liegt 0_V in U , und für alle $v \in U$ gilt $u +_U 0_V = u + 0_V = u$ und $0_V +_U u = 0_V + u = u$. Damit besitzt 0_V die Eigenschaften des Neutralelements in $(U, +_U)$. Sei nun $v \in U$ vorgegeben. Nach Voraussetzung liegt der Vektor $-v = (-1)v$ in U . Außerdem gilt $v +_U (-v) = v + (-v) = 0_V$ und $(-v) +_U v = (-v) + v = 0_V$. Also besitzt jedes $v \in U$ in $(U, +_U)$ ein Inverses, nämlich $-v$. Insgesamt bedeutet dies, dass $(U, +_U)$ eine Gruppe ist. Das Kommutativgesetz erhält man durch die Rechnung $v +_U w = v + w = w + v = w +_U v$ für alle $v, w \in U$.

Nun müssen wir noch die Eigenschaften (ii) (a)-(d) aus Definition (6.1) überprüfen. Seien dazu $v, w \in U$ und $\lambda, \mu \in K$ vorgegeben. Es gilt $(\lambda + \mu) \cdot_U v = (\lambda + \mu) \cdot v = \lambda \cdot v + \mu \cdot v = \lambda \cdot_U v +_U \mu \cdot_U v$. Ebenso erhält man $\lambda \cdot_U (v +_U w) = \lambda \cdot (v + w) = \lambda \cdot v + \lambda \cdot w = \lambda \cdot_U v +_U \lambda \cdot_U w$. Weiter gilt $(\lambda \mu) \cdot_U v = (\lambda \mu) \cdot v = \lambda \cdot (\mu \cdot v) = \lambda \cdot_U (\mu \cdot_U v)$ und schließlich $1_K \cdot_U v = 1_K \cdot v = v$. $\square$

Man sieht, dass das Nachrechnen der Vektorraum-Axiome in $(U, +_U, \cdot_U)$ eine ziemlich Routineangelegenheit war: Überall wurden nur die Symbole $+_U$ und $\cdot_U$ durch $+$ und $\cdot$ ersetzt und anschließend verwendet, dass die Axiome im Vektorraum V gültig sind.

Folgende konkrete Beispiele lassen sich für Untervektorräume angeben.

- (i) Ist V ein beliebiger K -Vektorraum, dann sind $\{0_V\}$ und V Untervektorräume von V .
- (ii) Für jedes $v \in V$ ist $\langle v \rangle_K = \{\lambda v \mid \lambda \in K\}$ ein Untervektorraum. Im Fall $v \neq 0_V$ bezeichnet man ihn als **lineare Gerade**. Für beliebige v, w ist auch durch

$$\langle v, w \rangle_K = \{\lambda v + \mu w \mid \lambda, \mu \in K\}$$

ein Untervektorraum gegeben. Ist $v \neq 0_V$ und $w \notin \langle v \rangle_K$ (oder äquivalent, $v \notin \langle w \rangle_K$ und $w \notin \langle v \rangle_K$), dann nennt man $\langle v, w \rangle_K$ eine **lineare Ebene**.

(6.6) Definition Eine Teilmenge $A \subseteq V$ wird **affiner Unterraum** von V genannt, wenn entweder $A = \emptyset$ gilt oder ein Untervektorraum U und ein Vektor $v \in V$ existieren, so dass

$$A = v + U = \{v + u \mid u \in U\} \quad \text{erfüllt ist.}$$

Betrachten wir einige konkrete Beispiele für affine Unterräume.

- (i) Seien $u, v \in V$. Dann ist $u + \langle v \rangle_K = \{u + \lambda v \mid \lambda \in K\}$ ein affiner Unterraum. Im Fall $v \neq 0_V$ bezeichnet man ihn als **affine Gerade**.
- (ii) Für beliebige $u, v, w \in V$ ist durch $u + \langle v, w \rangle_K = \{u + \lambda v + \mu w \mid \lambda, \mu \in K\}$ ein affiner Unterraum gegeben. Ist $v \neq 0_V$ und w kein skalares Vielfaches von v (also $w \neq \lambda v$ für alle $\lambda \in K$), dann nennt man $u + \langle v, w \rangle_K$ eine **affine Ebene**.

(6.7) Proposition Sei V ein K -Vektorraum und $\emptyset \neq A \subseteq V$ ein affiner Unterraum.

- (i) Es gibt *genau einen* Untervektorraum U von V , so dass die Gleichung $A = v + U$ für ein $v \in V$ erfüllt ist.
- (ii) Für jeden Vektor $w \in A$ erfüllt der Untervektorraum U aus Teil (i) die Gleichung $A = w + U$.

Wir nennen U den **zu A gehörenden Untervektorraum** und bezeichnen ihn mit $\mathcal{L}(A)$.

Beweis: zu (i) Nehmen wir an, dass $v, v' \in V$ Vektoren und U, U' Untervektorräume von V mit $v + U = A = v' + U'$ sind. Wegen $v \in A$ gilt $v = v' + u_0$ für ein $u_0 \in U'$, und wegen $v' \in A$ gilt $v' = v + u_1$ für ein $u_1 \in U$. Der Differenzvektor $v' - v$ ist also sowohl in U als auch in U' enthalten. Wir beweisen nun die Gleichung $U = U'$.

„ $\subseteq$ “ Ist $u \in U$, dann liegt $v + u$ in A , und folglich gibt es ein $u' \in U'$ mit $v + u = v' + u'$. Es folgt $u = (v' - v) + u' \in U'$. „ $\supseteq$ “ Ist $u' \in U'$ vorgegeben, dann gilt $v' + u'$ in A , es gibt also ein $u \in U$ mit $v' + u' = v + u$. Daraus folgt $u' = (v - v') + u \in U$.

zu (ii) Sei $U = \mathcal{L}(A)$, $v \in V$ ein Vektor mit $A = v + U$ und $w \in A$ ein beliebiges Element. Dann gibt es ein $u \in U$ mit $w = v + u$. Wir beweisen nun die Gleichung $v + U = w + U$. „ $\subseteq$ “ Ist $v_1 \in v + U$, dann gibt es ein $u_1 \in U$ mit $v_1 = v + u_1$, und es folgt $v_1 = (w - u) + u_1 = w + (u_1 - u) \in w + U$. „ $\supseteq$ “ Ist $w_1 \in w + U$, dann existiert ein $u_1 \in U$ mit $w_1 = w + u_1$. Es folgt $w_1 = w + u_1 = (v + u) + u_1 = v + (u + u_1) \in v + U$. $\square$

(6.8) Definition Seien $(V, +_V, \cdot_V)$ und $(W, +_W, \cdot_W)$ K -Vektorräume. Eine Abbildung $\phi : V \rightarrow W$ heißt **K -lineare Abbildung** oder **Homomorphismus** von K -Vektorräumen, wenn folgende Bedingungen erfüllt sind.

- (i) $\phi(v +_V w) = \phi(v) +_W \phi(w)$ für alle $v, w \in V$
- (ii) $\phi(\lambda \cdot_V v) = \lambda \cdot_W \phi(v)$ für alle $v \in V$ und $\lambda \in K$

Wenn aus dem Zusammenhang heraus klar ist, über welchem Körper die Vektorräume V und W definiert sind, wird statt von einer K -linearen auch einfach von einer linearen Abbildung gesprochen.

(6.9) Lemma Ist $\phi : V \rightarrow W$ eine lineare Abbildung. Dann gilt $\phi(0_V) = 0_W$, $\phi(-v) = -\phi(v)$ und $\phi(v - w) = \phi(v) - \phi(w)$ für alle $v, w \in V$.

Beweis: Die erste Gleichung erhält man mit Hilfe der Eigenschaft (ii) von linearen Abbildungen durch $\phi(0_V) = \phi(0_K \cdot_V 0_V) = 0_K \cdot_W \phi(0_V) = 0_W$. Die zweite ergibt sich durch die Rechnung

$$\phi(-v) = \phi((-1_K) \cdot_V v) = (-1_K) \cdot_W \phi(v) = -\phi(v).$$

Die dritte Gleichung schließlich erhält man durch

$$\phi(v - w) = \phi(v +_V (-w)) = \phi(v) +_W \phi(-w) = \phi(v) +_W (-\phi(w)) = \phi(v) - \phi(w). \quad \square$$

Sei V ein K -Vektorraum, $n \in \mathbb{N}$, und seien $v_1, \dots, v_n \in V$ beliebige Vektoren. Wir verwenden den Ausdruck $\sum_{k=1}^n v_k$ als Kurzschreibweise für die Summe $v_1 + \dots + v_n$ in V .

(6.10) Lemma Seien V, W K -Vektorräume, $n \in \mathbb{N}$, außerdem $v_1, \dots, v_n \in V$ und $\phi : V \rightarrow W$ eine lineare Abbildung. Dann gilt

$$\phi\left(\sum_{k=1}^n v_k\right) = \sum_{k=1}^n \phi(v_k).$$

Beweis: Wir beweisen die Aussage durch vollständige Induktion über n . Im Fall $n = 1$ lautet die Behauptung nur $\phi(v_1) = \phi(v_1)$ für alle $v_1 \in V$ und ist offensichtlich erfüllt. Sei nun $n \in \mathbb{N}$, und setzen wir die Aussage für dieses n voraus. Seien $v_1, \dots, v_{n+1} \in V$ beliebige Vektoren. Dann erhalten wir

$$\begin{aligned} \phi\left(\sum_{k=1}^{n+1} v_k\right) &= \phi\left(\sum_{k=1}^n v_k + v_{n+1}\right) = \phi\left(\sum_{k=1}^n v_k\right) + \phi(v_{n+1}) \stackrel{(*)}{=} \\ &\sum_{k=1}^n \phi(v_k) + \phi(v_{n+1}) = \sum_{k=1}^{n+1} \phi(v_k), \end{aligned}$$

wobei an der Stelle $(*)$ die Induktionsvoraussetzung angewendet wurde. $\square$

Wir haben in § 5 angemerkt, dass jedes $v \in K^n$ auf natürliche Weise mit einer Matrix in $\mathcal{M}_{1 \times n, K}$ und auch mit einer Matrix in $\mathcal{M}_{n \times 1, K}$ identifiziert werden kann. Ist nun $A \in \mathcal{M}_{n, K}$ und betrachten wir $v \in K^n$ gemäß der zweiten Möglichkeit als Matrix mit einer Spalte und n Einträgen, dann können wir das Matrixprodukt $Av \in K^m$ bilden. Man bezeichnet den Vektor $w = Av$ dann als **Matrix-Vektor-Produkt** von A und v . Aus der Definition des Matrixprodukts ergibt sich, dass die Komponenten des Vektors w durch $w_i = \sum_{j=1}^n a_{ij} v_j$ gegeben sind, für $1 \leq i \leq m$.

(6.11) Proposition Seien $m, n \in \mathbb{N}$ und $A \in \mathcal{M}_{m \times n, K}$. Dann ist durch $\phi_A : K^n \rightarrow K^m$, $v \mapsto Av$ eine lineare Abbildung gegeben.

Beweis: Seien $v, w \in K^n$ und $\lambda \in K$ vorgegeben. Auf Grund der Rechenregeln für das Matrixprodukt aus Teil (i) von Proposition (5.15) gilt

$$\phi_A(v + w) = A(v + w) = Av + Aw = \phi_A(v) + \phi_A(w)$$

und $\phi_A(\lambda v) = A(\lambda v) = \lambda Av = \lambda \phi_A(v)$. $\square$

(6.12) Proposition Seien U, V, W drei K -Vektorräume und $\phi : U \rightarrow V$, $\psi : V \rightarrow W$ lineare Abbildungen. Dann ist $\psi \circ \phi$ eine lineare Abbildung von U nach W . Ist ϕ bijektiv, dann ist ϕ^{-1} eine lineare Abbildung von V nach U .

Beweis: Wir überprüfen die Linearität der Abbildung $\psi \circ \phi$. Seien dazu $v, w \in U$ und $\lambda \in K$ vorgegeben. Dann gilt

$$\begin{aligned} (\psi \circ \phi)(v +_U w) &= \psi(\phi(v +_U w)) = \psi(\phi(v) +_V \phi(w)) = \psi(\phi(v)) +_W \psi(\phi(w)) \\ &= (\psi \circ \phi)(v) +_W (\psi \circ \phi)(w) \end{aligned}$$

und $(\psi \circ \phi)(\lambda v) = \psi(\phi(\lambda v)) = \psi(\lambda \phi(v)) = \lambda \psi(\phi(v)) = \lambda (\psi \circ \phi)(v)$. Setzen wir nun voraus, dass ϕ bijektiv ist. Um zu zeigen, dass die Abbildung ϕ^{-1} linear ist, seien $v, w \in V$ und $\lambda \in K$ vorgegeben. Sei $v' = \phi^{-1}(v)$ und $w' = \phi^{-1}(w)$. Unter Verwendung der Linearität von ϕ erhalten wir

$$\begin{aligned} \phi^{-1}(v) +_U \phi^{-1}(w) &= v' +_U w' = \text{id}_U(v' +_U w') = (\phi^{-1} \circ \phi)(v' +_U w') \\ &= \phi^{-1}(\phi(v' +_U w')) = \phi^{-1}(\phi(v') +_V \phi(w')) = \phi^{-1}(v +_V w). \end{aligned}$$

Ebenso gilt $\lambda \phi^{-1}(v) = \lambda v' = \text{id}_U(\lambda v') = (\phi^{-1} \circ \phi)(\lambda v') = \phi^{-1}(\phi(\lambda v')) = \phi^{-1}(\lambda \phi(v')) = \phi^{-1}(\lambda v)$. □

(6.13) Definition Eine lineare Abbildung $\phi : V \rightarrow W$ heißt

- (i) **Monomorphismus** (von K -Vektorräumen), wenn ϕ injektiv ist,
- (ii) **Epimorphismus**, wenn ϕ surjektiv ist,
- (iii) **Isomorphismus**, wenn ϕ bijektiv ist.

Eine lineare Abbildung $\phi : V \rightarrow V$ bezeichnet man als **Endomorphismus** von V , und ist sie außerdem bijektiv, dann spricht man von einem **Automorphismus**. Zwei K -Vektorräume V, W werden **isomorph** genannt, wenn ein Isomorphismus $\phi : V \rightarrow W$ existiert.

(6.14) Folgerung Die Menge der Automorphismen eines K -Vektorraums V ist mit der Komposition $\circ$ von Abbildungen als Verknüpfung eine Gruppe. Man bezeichnet sie mit $\text{GL}(V)$ und nennt sie die **allgemeine lineare Gruppe** des Vektorraums V .

Beweis: Proposition (6.12) zeigt, dass für gegebene Automorphismen ϕ, ψ des K -Vektorraums V auch die Abbildungen $\psi \circ \phi$ und ϕ^{-1} Automorphismen von V sind. Die Assoziativität ergibt sich aus der allgemeinen Regel $h \circ (g \circ f) = (h \circ g) \circ f$ für beliebige Abbildungen zwischen Mengen. Die identische Abbildung id_V besitzt die definierende Eigenschaft des Neutralelements (es gilt $\phi \circ \text{id}_V = \text{id}_V \circ \phi = \phi$ für jeden Automorphismus ϕ von V), und die Umkehrabbildung ϕ^{-1} von ϕ erfüllt die Bedingung $\phi^{-1} \circ \phi = \phi \circ \phi^{-1} = \text{id}_V$ für das inverse Element. □

(6.15) Definition Eine Abbildung $\psi : V \rightarrow W$ zwischen zwei K -Vektorräumen wird **affin-lineare** Abbildung genannt, wenn eine lineare Abbildung $\phi : V \rightarrow W$ und ein Vektor $w \in W$ existieren, so dass $\psi(v) = w + \phi(v)$ für alle $v \in V$ erfüllt ist. Eine bijektive affin-lineare Abbildung $\psi : V \rightarrow V$ wird **Affinität** von V genannt.

(6.16) Proposition Seien V, W K -Vektorräume und $\phi : V \rightarrow W$ eine lineare Abbildung. Ferner seien $V' \subseteq V$ und $W' \subseteq W$ Untervektorräume. Dann sind die Teilmengen

$$\phi(V') = \{\phi(v) \mid v \in V'\} \quad \text{und} \quad \phi^{-1}(W') = \{v \in V \mid \phi(v) \in W'\}$$

Untervektorräume von W bzw. von V .

Beweis: Wir rechnen die Untervektorraum-Axiome für beide Teilmengen direkt nach. Seien $w, w' \in \phi(V')$ und $\lambda \in K$. Dann gibt es nach Definition von $\phi(V')$ Vektoren $v, v' \in V'$ mit $w = \phi(v)$ und $w' = \phi(v')$. Da V' ein Untervektorraum ist, gilt $v + v' \in V'$ und damit

$$w + w' = \phi(v) + \phi(v') = \phi(v + v') \in \phi(V').$$

Ebenso gilt $\lambda v \in V'$ auf Grund der Untervektorraum-Eigenschaft und somit $\lambda w = \lambda \phi(v) = \phi(\lambda v) \in \phi(V')$.

Nun zeigen wir, dass auch $\phi^{-1}(W')$ ein Untervektorraum ist. Seien dazu $v, v' \in \phi^{-1}(W')$ und $\lambda \in K$ vorgegeben. Dann gilt $\phi(v), \phi(v') \in W'$ und $\phi(v) + \phi(v') \in W'$, da W' ein Untervektorraum von W ist. Aus $\phi(v + v') = \phi(v) + \phi(v') \in W'$ folgt $v + v' \in \phi^{-1}(W')$. Da auch $\lambda \phi(v)$ in W' liegt, erhalten wir $\phi(\lambda v) = \lambda \phi(v) \in W'$ und somit $\lambda v \in \phi^{-1}(W')$. $\square$

(6.17) Definition Seien V, W zwei K -Vektorräume und $\phi : V \rightarrow W$ eine lineare Abbildung. Dann nennt man

- (i) $\ker(\phi) = \phi^{-1}(\{0_W\}) = \{v \in V \mid \phi(v) = 0_W\}$ den **Kern** und
- (ii) $\text{im}(\phi) = \phi(V) = \{\phi(v) \mid v \in V\}$ das **Bild** von ϕ .

Nach Prop. (6.16) ist $\ker(\phi)$ ein Untervektorraum von V und $\text{im}(\phi)$ ein Untervektorraum von W .

(6.18) Proposition Seien V, W K -Vektorräume und $\phi : V \rightarrow W$ eine lineare Abbildung.

- (i) Die Abbildung ϕ ist genau dann surjektiv, wenn $\text{im}(\phi) = W$ gilt.
- (ii) Sie ist genau dann injektiv, wenn $\ker(\phi) = \{0_V\}$ erfüllt ist.

Beweis: Aussage (i) ist nach Definition der Surjektivität unmittelbar klar. Zum Beweis von (ii) setzen wir zunächst voraus, dass ϕ injektiv ist. Die Inklusion $\{0_V\} \subseteq \ker(\phi)$ ist erfüllt, weil der Kern ein Untervektorraum von V ist. Zum Nachweis von $\ker(\phi) \subseteq \{0_V\}$ sei $v \in \ker(\phi)$ vorgegeben. Dann gilt $\phi(v) = 0_W = \phi(0_V)$, und aus der Injektivität von ϕ folgt $v = 0_V$.

Setzen wir nun umgekehrt die Gleichung $\ker(\phi) = \{0_V\}$ voraus, und beweisen wir die Injektivität von ϕ . Seien dazu $v, v' \in V$ mit $\phi(v) = \phi(v')$ vorgegeben. Dann gilt $\phi(v' - v) = \phi(v') - \phi(v) = 0_W$ und somit $v' - v \in \ker(\phi)$. Aus der Voraussetzung an den Kern folgt $v' - v = 0_V \Leftrightarrow v = v'$. $\square$

(6.19) Definition Sei K ein Körper, und seien $m, n \in \mathbb{N}$. Ein **lineares Gleichungssystem** über K bestehend aus m Gleichungen in n Unbekannten $x_1, \dots, x_n$ ist ein System von Gleichungen der Form

$$\begin{array}{ccccccc} a_{11}x_1 & + & a_{12}x_2 & + & \cdots & + & a_{1n}x_n & = & b_1 \\ \vdots & & \vdots & & & & \vdots & = & \vdots \\ a_{m1}x_1 & + & a_{m2}x_2 & + & \cdots & + & a_{mn}x_n & = & b_m \end{array}$$

wobei $a_{ij} \in K$ und $b_i \in K$ für $1 \leq i \leq m$ und $1 \leq j \leq n$ ist. Gilt $b_i = 0$ für $1 \leq i \leq m$, dann spricht man von einem **homogenen**, ansonsten von einem **inhomogenen** LGS.

Mit Hilfe der Matrixschreibweise lässt sich ein lineares Gleichungssystem in deutlich kompakterer Form darstellen.

(6.20) Definition Die Matrix $A = (a_{ij})_{m \times n, K}$ in Definition (6.19) wird **Koeffizientenmatrix** des LGS genannt. Die Matrix $\tilde{A} = (\tilde{a}_{ij})_{m \times (n+1), K}$ gegeben durch $\tilde{a}_{ij} = a_{ij}$ für $1 \leq i \leq m$, $1 \leq j \leq n$ und $\tilde{a}_{i, n+1} = b_i$ für $1 \leq i \leq m$ heißt **erweiterte Koeffizientenmatrix**. Bezeichnet x die $n \times 1$ -Matrix mit den Einträgen $x_1, \dots, x_n$, dann kann das lineare Gleichungssystem in der Form $Ax = b$ dargestellt werden. Die Menge

$$\mathcal{L} = \mathcal{L}_{A,b} = \{c \in K^n \mid Ac = b\}$$

bezeichnet man als **Lösungsmenge**, deren Elemente als **Lösungen** des linearen Gleichungssystems. Ist $\mathcal{L}_{A,b} \neq \emptyset$, dann bezeichnet man $Ax = b$ als **lösbar**. Besteht $\mathcal{L}_{A,b}$ aus einem einzigen Element, dann spricht man von **eindeutiger Lösbarkeit**.

Wir zwei konkrete Beispiele über dem Körper $K = \mathbb{R}$ der reellen Zahlen. Das lineare Gleichungssystem

$$\begin{array}{rrcr} 3x & & + & 2z & = & 8 \\ -x & + & 2y & + & 5z & = & -3 \\ & & 7y & - & 2z & = & -23 \end{array}$$

hat in Matrixschreibweise die Form $Ax = b$ mit

$$A = \begin{pmatrix} 3 & 0 & 2 \\ -1 & 2 & 5 \\ 0 & 7 & -2 \end{pmatrix} \quad \text{und} \quad b = \begin{pmatrix} 8 \\ -3 \\ -23 \end{pmatrix}.$$

Die erweiterte Koeffizientenmatrix des Systems ist also gegeben durch

$$\tilde{A} = \begin{pmatrix} 3 & 0 & 2 & 8 \\ -1 & 2 & 5 & -3 \\ 0 & 7 & -2 & -23 \end{pmatrix}.$$

Mit den Methoden, die im nächsten Kapitel beschrieben werden, kann man ausrechnen, dass die Lösungsmenge des Systems durch $\mathcal{L}_{A,b} = \{(2, -3, 1)\}$ gegeben ist. Das System ist also eindeutig lösbar. Betrachtet man dagegen das lineare Gleichungssystem

$$\begin{aligned} 3x + 5y - 3z &= -4 \\ 2x - 8y + 7z &= 11 \\ 5x - 3y + 4z &= 7 \end{aligned}$$

dann ist die Lösungsmenge gegeben durch

$$\mathcal{L} = \left\{ \begin{pmatrix} 1 \\ -2 \\ -1 \end{pmatrix} + \lambda \cdot \begin{pmatrix} 11 \\ -27 \\ -34 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\}.$$

In diesem Fall liegt also Lösbarkeit, aber keine eindeutige Lösbarkeit vor.

(6.21) Satz Seien K ein Körper, $m, n \in \mathbb{N}$, $A \in \mathcal{M}_{m \times n, K}$ und $b \in K^m$. Dann ist $\mathcal{L}_{A,b}^{\text{hom}} = \mathcal{L}_{A,0}$ ein Untervektorraum des K^n . Ist das lineare Gleichungssystem $Ax = b$ lösbar und ist $c \in K^n$ eine Lösung, dann gilt $\mathcal{L}_{A,b} = c + \mathcal{L}_{A,b}^{\text{hom}}$. Die Lösungsmenge $\mathcal{L}_{A,b} \subseteq K^n$ ist also ein affiner Unterraum des K^n .

Beweis: Nach Proposition (6.11) ist $K^n \rightarrow K^m$, $v \mapsto Av$ eine lineare Abbildung, und $\mathcal{L}_{A,b}^{\text{hom}}$ ist der Kern dieser linearen Abbildung. Als Kern einer linearen Abbildung ist $\mathcal{L}_{A,b}^{\text{hom}}$ ein Untervektorraum im K^n . Für den Beweis der Gleichung $\mathcal{L}_{A,b} = c + \mathcal{L}_{A,b}^{\text{hom}}$ bemerken wir zunächst, dass $Ac = b$ gilt, weil c eine Lösung des Systems $Ax = b$ ist. Um die Inklusion „ $\supseteq$ “ nachzuweisen, sei $v \in c + \mathcal{L}_{A,b}^{\text{hom}}$, also $v = c + w$ für ein $w \in \mathcal{L}_{A,b}^{\text{hom}}$. Es gilt dann $Aw = 0_{K^m}$, und wir erhalten $Av = A(c + w) = Ac + Aw = b + 0_{K^m} = b$. Dies zeigt, dass v in $\mathcal{L}_{A,b}$ enthalten ist. Für den Nachweis von „ $\subseteq$ “ sei $v \in \mathcal{L}_{A,b}$. Dann gilt $Av = b$. Setzen wir $w = v - c$, dann folgt $Aw = A(v - c) = Av - Ac = b - b = 0_{K^m}$, also $w \in \mathcal{L}_{A,b}^{\text{hom}}$. Es folgt $v = c + w \in c + \mathcal{L}_{A,b}^{\text{hom}}$. $\square$

Für die Anwendungen in den späteren Kapiteln führen wir noch ein weiteres wichtiges Beispiel für einen K -Vektorraum ein, den Vektorraum der linearen Abbildungen.

Sei K ein Körper, X eine Menge und V ein K -Vektorraum. Weiter sei $\text{Abb}(X, V)$ die Menge der Abbildungen $X \rightarrow V$. Wir definieren auf $\text{Abb}(X, V)$ eine Verknüpfung $+$, in dem wir $\varphi, \psi \in \text{Abb}(X, V)$ das Element $\varphi + \psi \in \text{Abb}(X, V)$ gegeben durch $(\varphi + \psi)(x) = \varphi(x) + \psi(x)$ für alle $x \in X$ zuordnen. Außerdem definieren wir eine Abbildung

$$\cdot : K \times \text{Abb}(X, V) \rightarrow \text{Abb}(X, V),$$

indem wir für $\lambda \in K$ und $\varphi \in \text{Abb}(X, V)$ die das Element $\lambda \cdot \varphi$ gegeben durch $(\lambda\varphi)(x) = \lambda\varphi(x)$ für alle $x \in X$ definieren.

(6.22) Proposition Das Tripel $(\text{Abb}(X, V), +, \cdot)$ ist ein K -Vektorraum.

Beweis: Wir müssen überprüfen, dass $(\text{Abb}(X, V), +, \cdot)$ die in Definition (6.1) aufgezählten Eigenschaften besitzt. Zunächst überprüfen wir, dass es sich bei $(\text{Abb}(X, V), +)$ um eine abelsche Gruppe handelt. Die Verknüpfung $+$ auf $\text{Abb}(X, V)$, denn für alle $\varphi, \psi, \phi \in \text{Abb}(X, V)$ und alle $x \in X$ gilt

$$\begin{aligned} ((\varphi + \psi) + \phi)(x) &= (\varphi + \psi)(x) + \phi(x) = (\varphi(x) + \psi(x)) + \phi(x) = \varphi(x) + (\psi(x) + \phi(x)) \\ &= \varphi(x) + (\psi + \phi)(x) = (\varphi + (\psi + \phi))(x) \end{aligned}$$

also $(\varphi + \psi) + \phi = \varphi + (\psi + \phi)$. Ebenso gilt $(\varphi + \psi)(x) = \varphi(x) + \psi(x) = \psi(x) + \varphi(x) = (\psi + \varphi)(x)$ und somit $\varphi + \psi = \psi + \varphi$. Das Paar $(\text{Abb}(X, V), +)$ ist somit eine abelsche Halbgruppe. Definieren wir $0_{\text{Abb}(X, V)}$ durch $0_{\text{Abb}(X, V)}(x) = 0_V$ für alle $x \in X$, dann gilt für alle $\varphi \in \text{Abb}(X, V)$ und alle $x \in X$ jeweils $(0_{\text{Abb}(X, V)} + \varphi)(x) = 0_{\text{Abb}(X, V)}(x) + \varphi(x) = 0_V + \varphi(x) = \varphi(x)$, also $0_{\text{Abb}(X, V)} + \varphi = \varphi$. Auf Grund des bereits bewiesenen Kommutativgesetzes folgt daraus auch $\varphi + 0_{\text{Abb}(X, V)} = \varphi$. Damit ist gezeigt, dass $0_{\text{Abb}(X, V)}$ in der Halbgruppe $(\text{Abb}(X, V), +)$ ein Neutralelement ist, es sich also um ein Monoid handelt.

Sei schließlich $\varphi \in \text{Abb}(X, V)$ vorgegeben und $-\varphi \in \text{Abb}(X, V)$ gegeben durch $(-\varphi)(x) = -\varphi(x)$ für alle $x \in X$. Für jedes $x \in X$ gilt dann $((-\varphi) + \varphi)(x) = (-\varphi)(x) + \varphi(x) = (-\varphi(x)) + \varphi(x) = 0_V = 0_{\text{Abb}(X, V)}(x)$, also $\varphi + (-\varphi) = 0_{\text{Abb}(X, V)}$. Auf Grund des Kommutativgesetzes gilt auch $(-\varphi) + \varphi = 0_{\text{Abb}(X, V)}$. Jedes Element im Monoid $(\text{Abb}(X, V), +)$ ist also invertierbar, somit ist $(\text{Abb}(X, V), +)$ insgesamt eine abelsche Gruppe.

Nun überprüfen wir noch die übrigen in Definition (6.1) genannten Rechenregeln. Seien dazu $\lambda, \mu \in K$ und $\varphi, \psi \in \text{Abb}(X, V)$ vorgegeben. Für jedes $x \in X$ gelten die Gleichungen

$$((\lambda + \mu) \cdot \varphi)(x) = (\lambda + \mu)\varphi(x) = \lambda\varphi(x) + \mu\varphi(x) = (\lambda \cdot \varphi)(x) + (\mu \cdot \varphi)(x) = ((\lambda \cdot \varphi) + (\mu \cdot \varphi))(x)$$

und

$$\begin{aligned} (\lambda \cdot (\varphi + \psi))(x) &= \lambda(\varphi + \psi)(x) = \lambda(\varphi(x) + \psi(x)) = \lambda\varphi(x) + \lambda\psi(x) = \\ &= (\lambda \cdot \varphi)(x) + (\lambda \cdot \psi)(x) = (\lambda \cdot \varphi + \lambda \cdot \psi)(x) \end{aligned}$$

ebenso

$$((\lambda\mu) \cdot \varphi)(x) = (\lambda\mu)\varphi(x) = \lambda(\mu\varphi(x)) = \lambda((\mu \cdot \varphi)(x)) = (\lambda \cdot (\mu \cdot \varphi))(x)$$

und schließlich $(1_K \cdot \varphi)(x) = 1_K\varphi(x) = \varphi(x)$. Es ist also $(\lambda + \mu) \cdot \varphi = \lambda \cdot \varphi + \mu \cdot \varphi$, $\lambda \cdot (\varphi + \psi) = \lambda \cdot \varphi + \lambda \cdot \psi$, $(\lambda\mu) \cdot \varphi = \lambda \cdot (\mu \cdot \varphi)$ und $1_K \cdot \varphi = \varphi$. $\square$

(6.23) Satz Sei K ein Körper, und seien V und W zwei K -Vektorräume. Wir bezeichnen mit $\text{Hom}_K(V, W)$ die Menge der linearen Abbildungen $V \rightarrow W$. Für vorgegebene $\varphi, \psi \in \text{Hom}_K(V, W)$ und $\lambda \in K$ seien die Abbildungen $\varphi + \psi : V \rightarrow W$ und $\lambda \cdot \varphi : V \rightarrow W$ definiert durch $(\varphi + \psi)(v) = \varphi(v) + \psi(v)$ und $(\lambda \cdot \varphi)(v) = \lambda\varphi(v)$ für alle $v \in V$. Dann ist $(\text{Hom}_K(V, W), +, \cdot)$ ein K -Vektorraum.

Beweis: Nach Proposition (6.22) ist $\text{Abb}(V, W)$ ein K -Vektorraum. Nach Satz (6.5) genügt es zu zeigen, dass $\text{Hom}_K(V, W)$ ein Untervektorraum dieses K -Vektorraums ist. Wie wir im Beweis der Proposition gesehen haben, ist der Nullvektor von $\text{Abb}(V, W)$ gegeben durch $0_{\text{Abb}(V, W)}(v) = 0_W$ für alle $v \in V$. Es handelt sich dabei um eine lineare Abbildung $V \rightarrow W$, also um ein Element von $\text{Hom}_K(V, W)$, denn für alle $v, v' \in V$ und alle $\alpha \in K$ gilt $0_{\text{Abb}(V, W)}(v + v') = 0_W = 0_W + 0_W = 0_{\text{Abb}(V, W)}(v) + 0_{\text{Abb}(V, W)}(v')$ und $0_{\text{Abb}(V, W)}(\alpha v) = 0_W = \alpha 0_W = \alpha 0_{\text{Abb}(V, W)}(v)$.

Seien nun $\varphi, \psi \in \text{Hom}_K(V, W)$ und $\lambda \in K$ vorgegeben. Zu zeigen ist $\varphi + \psi \in \text{Hom}_K(V, W)$ und $\lambda \cdot \varphi \in \text{Hom}_K(V, W)$. Für den Nachweis seien $v, v' \in V$ und $\alpha \in K$ vorgegeben. Es gilt $(\varphi + \psi)(\alpha v) = \varphi(\alpha v) + \psi(\alpha v) = \alpha \varphi(v) + \alpha \psi(v) = \alpha(\varphi(v) + \psi(v)) = \alpha(\varphi + \psi)(v)$, wobei im zweiten Schritt verwendet wurde, dass φ und ψ lineare Abbildungen sind. Die Verträglichkeit von $\varphi + \psi$ mit der Vektoraddition erhält man durch die Rechnung

$$\begin{aligned} (\varphi + \psi)(v + v') &= \varphi(v + v') + \psi(v + v') = \varphi(v) + \varphi(v') + \psi(v) + \psi(v') = \\ &= \varphi(v) + \psi(v) + \varphi(v') + \psi(v') = (\varphi + \psi)(v) + (\varphi + \psi)(v'). \end{aligned}$$

Damit ist insgesamt $\varphi + \psi \in \text{Hom}_K(V, W)$ nachgewiesen. Ebenso erhält man $(\lambda \cdot \varphi)(v + v') = \lambda \varphi(v + v') = \lambda(\varphi(v) + \varphi(v')) = \lambda \varphi(v) + \lambda \varphi(v') = (\lambda \cdot \varphi)(v) + (\lambda \cdot \varphi)(v')$ und $(\lambda \cdot \varphi)(\alpha v) = \lambda \varphi(\alpha v) = \lambda \alpha \varphi(v) = \alpha \lambda \varphi(v) = \alpha(\lambda \cdot \varphi)(v)$. Also gilt auch $\lambda \cdot \varphi \in \text{Hom}_K(V, W)$. $\square$

§ 7. Die Lösung linearer Gleichungssysteme

Inhaltsübersicht

In diesem Kapitel behandeln wir ein allgemeines Verfahren, dass die Bestimmung der Lösungsmenge von beliebig großen LGS ermöglicht. Dabei gehen wir davon aus, dass das LGS durch seine erweiterte Koeffizientenmatrix (siehe § 6) gegeben ist. Diese Matrix bringt man durch eine fest vorgegebene Folge von Umformungsschritten auf eine sog. normierte Zeilenstufenform. Die Lösungsmenge kann dann auf einfache Weise von der Matrix abgelesen werden, egal ob diese leer oder einelementig ist, oder aus mehreren Elementen besteht. Neben dem Gaußschen Eliminationsverfahren beschäftigen wir uns mit Rechenoperationen für Matrizen und untersuchen, wie diese mit den elementaren Umformungen eines LGS aus dem ersten Kapitel zusammenhängen.

Wichtige Begriffe und Sätze

- (normierte) Zeilenstufenform (ZSF)
- Kennzahlen der ZSF, Zeilenköpfe, Rang einer Matrix in ZSF
- Einheitsvektoren
- Elementarmatrix, elementare Zeilen- bzw. Spaltenumformung
- Blockschreibweise für Matrizen
- Gauß'sches Eliminationsverfahren
- Invertierbarkeitskriterium für Matrizen

(7.1) Definition Eine Matrix $A \in \mathcal{M}_{m \times n, K}$ befindet sich in **Zeilenstufenform** (kurz ZSF), wenn $A = \mathbf{0}^{(m \times n)}$ gilt oder folgende Bedingung erfüllt ist: Es gibt ein $r \in \{1, \dots, m\}$ und $j_1, \dots, j_r \in \{1, \dots, n\}$ mit $j_1 < j_2 < \dots < j_r$, so dass

- (i) $a_{ij_i} \neq 0_K$ für $1 \leq i \leq r$ und
- (ii) $a_{ij} = 0_K$ für $j < j_i$ oder $i > r$

erfüllt ist. Man nennt r den **Zeilenrang** einer solchen Matrix. Das Tupel $(r, j_1, \dots, j_r)$ bezeichnen wir insgesamt als die **Kennzahlen** der ZSF.

Die Positionen (i, j_i) mit $1 \leq i \leq r$ in der Matrix werden **Zeilenköpfe** genannt. Die Bedingung (i) besagt, dass die Einträge in den Zeilenköpfen ungleich Null sind. Nach Bedingung (ii) befinden sich links von den Zeilenköpfen nur Nulleinträge; in den „kopfloosen“ Zeilen sind alle Einträge gleich Null. Der Zeilenrang kann offenbar nie größer als $\min\{m, n\}$ werden.

(7.2) Definition Eine Matrix $A \in \mathcal{M}_{m \times n, K}$ befindet sich in **normierter** ZSF, wenn $A = \mathbf{0}^{(m \times n)}$ gilt oder wenn sie in ZSF mit den Kennzahlen $(r, j_1, \dots, j_r)$ vorliegt und außerdem die folgenden Bedingungen erfüllt sind: Es gilt $a_{ij_i} = 1_K$ für $1 \leq i \leq r$ und $a_{kj_i} = 0_K$ für $1 \leq i \leq r$ und $1 \leq k < i$.

Bei der normierten ZSF kommen also folgende Bedingungen hinzu: Die Einträge in den Zeilenköpfen sind gleich 1_K , und oberhalb der Zeilenköpfe befinden sich nur Nulleinträge. Bei einer Matrix A in normierter ZSF gilt also insgesamt $a_{ij} = 0$ in jedem der drei Fälle

$$(1) \quad i > r \quad (2) \quad i \leq r \text{ und } j < j_i \quad (3) \quad j = j_k \text{ für ein } k \in \{1, \dots, r\} \setminus \{i\};$$

in Worten, die Einträge der Matrix sind Null (1) unterhalb der r -ten Zeile, (2) links von den Spaltenköpfen und (3) in jeder Spalte, in der sich ein Zeilenkopf befindet, sind alle anderen Einträge gleich Null. Abgesehen davon können aber durchaus noch weitere Einträge von A gleich Null sein. Wir bemerken außerdem, dass eine Matrix $A \in \mathcal{M}_{m \times n, K}$ in normierter ZSF mit Zeilenrang $r = n$ in den oberen n Zeilen mit der Einheitsmatrix $E^{(n)}$ übereinstimmt, denn in diesem Fall muss $j_k = k$ für $1 \leq k \leq n$ gelten.

Wir geben einige konkrete Beispiele für Matrizen in Zeilenstufenform an.

$$A = \begin{pmatrix} 0 & 2 & 3 & 4 & 0 & 6 \\ 0 & 0 & 3 & 0 & 8 & 1 \\ 0 & 0 & 0 & 2 & 2 & 9 \\ 0 & 0 & 0 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Diese Matrix liegt in Zeilenstufenform vor, mit zugehörigen Kennzahlen $r = 4$, $j_1 = 2$, $j_2 = 3$, $j_3 = 4$ und $j_4 = 6$. Es besteht aber keine normierte Zeilenstufenform, denn beispielsweise sind die Einträge in den Zeilenköpfen $(1, 2)$, $(2, 3)$, $(3, 4)$ und $(4, 6)$ ungleich 1. Außerdem gibt es Einträge ungleich Null oberhalb der Zeilenköpfe.

$$B = \begin{pmatrix} 1 & 0 & 7 & 0 & 5 & 0 \\ 0 & 1 & 1 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 & 3 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Dies ist eine Matrix in normierter Zeilenstufenform, mit den Kennzahlen $r = 3$, $j_1 = 1$, $j_2 = 2$ und $j_3 = 4$. Auch die Einheitsmatrix

$$E^{(4)} = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

liegt in normierter Zeilenstufenform vor. Die Kennzahlen lauten $r = 4$ und $j_i = i$ für $1 \leq i \leq 4$.

Unser Ziel besteht darin, die Lösungsmenge eines linearen Gleichungssystems $Ax = b$ (mit $A \in \mathcal{M}_{n,K}$ und $b \in K^n$) zu bestimmen. Dabei konzentrieren wir uns zunächst auf den Fall, dass die erweiterte Koeffizientenmatrix $\tilde{A} = (Ab)$ in normierter Zeilenstufenform vorliegt. Nach Satz (6.21) genügt es, den Untervektorraum $\mathcal{L}_{A,b}^{\text{hom}} = \mathcal{L}_{A,0}$ von K^n und ein Element $c \in \mathcal{L}_{A,b}$ zu berechnen, denn dann gilt $\mathcal{L}_{A,b} = c + \mathcal{L}_{A,b}^{\text{hom}}$. Wir befassen uns als erstes mit der Bestimmung von $\mathcal{L}_{A,b}^{\text{hom}}$.

(7.3) Definition Für $1 \leq \ell \leq n$ bezeichnet $e_\ell \in K^n$ jeweils den ℓ -ten **Einheitsvektor** mit den Einträgen $e_{\ell j} = \delta_{\ell j}$.

In K^3 sind die drei Einheitsvektoren beispielsweise gegeben durch

$$e_1 = (1_K, 0_K, 0_K) \quad , \quad e_2 = (0_K, 1_K, 0_K) \quad \text{und} \quad e_3 = (0_K, 0_K, 1_K).$$

Sei nun A eine Matrix in normierter ZSF mit Kennzahlen $(r, j_1, \dots, j_r)$. Um die Lösungsmenge $\mathcal{L}_{A,0}$ von $Ax = 0_{K^m}$ anzugeben, definieren wir

$$S = \{1, \dots, n\} \setminus \{j_1, \dots, j_r\}$$

und definieren für jede Zahl $\ell \in S$ einen Vektor $v_\ell \in K^m$ durch

$$v_\ell = e_\ell - \sum_{k=1}^r a_{k\ell} e_{j_k}.$$

Der Vektor v_ℓ entsteht also aus dem Nullvektor dadurch, dass man die ℓ -te Komponente auf 1_K setzt und die Einträge $-a_{1\ell}, \dots, -a_{r\ell}$ der ℓ -ten Spalte auf die Positionen $j_1, \dots, j_r$ des Vektors verteilt. Es gilt also

$$v_{\ell j_k} = -a_{k\ell} \quad \text{für } 1 \leq k \leq r \quad \text{und} \quad v_{\ell j} = \delta_{\ell j} \quad \text{für alle } \ell \in S.$$

Mit Hilfe dieser Vektoren lässt sich nun die Lösungsmenge folgendermaßen darstellen.

(7.4) Satz Sei $\mathcal{L}^{\text{hom}} \subseteq K^n$ die Lösungsmenge eines homogenen LGS mit Koeffizientenmatrix A , und seien S und die Vektoren v_ℓ für $\ell \in S$ definiert wie oben.

(i) Im Fall $S = \emptyset$ gilt $\mathcal{L}^{\text{hom}} = \{0_{K^n}\}$.

(ii) Ist S nichtleer, dann ist die Lösungsmenge gegeben durch

$$\mathcal{L}^{\text{hom}} = \left\{ \sum_{\ell \in S} \lambda_\ell v_\ell \mid \lambda_\ell \in K \quad \forall \ell \in S \right\}.$$

Beweis: zu (i) Unter dieser Voraussetzung gilt $\{j_1, \dots, j_r\} = \{1, \dots, n\}$, woraus wiederum $r = n$ und somit $m \geq n$ folgt. Wie oben ausgeführt, stimmen bei Zeilenrang n die ersten n Zeilen von A mit der Einheitsmatrix $E^{(n)}$ überein. Es gilt also $a_{ij} = \delta_{ij}$ für $1 \leq i, j \leq n$ und $a_{ij} = 0_K$ falls $i > n$. Wir erinnern außerdem daran, dass nach Definition $\mathcal{L}^{\text{hom}} = \{w \in K^n \mid Aw = 0_{K^m}\}$ gilt. Für jeden Vektor $w \in K^n$ erhalten wir die Äquivalenz

$$\begin{aligned} w \in \mathcal{L}^{\text{hom}} &\Leftrightarrow Aw = 0_{K^m} \Leftrightarrow (Aw)_k = 0 \quad \text{für } 1 \leq k \leq m \Leftrightarrow \sum_{j=1}^n a_{kj} w_j = 0_K \quad \text{für } 1 \leq k \leq m \\ &\Leftrightarrow \sum_{j=1}^n \delta_{kj} w_j = 0_K \quad \text{für } 1 \leq k \leq n \Leftrightarrow w_k = 0_K \quad \text{für } 1 \leq k \leq n \Leftrightarrow w = 0_{K^n}. \end{aligned}$$

zu (ii) Hier beschränken wir uns auf den Nachweis, dass für jedes $\ell \in S$ der Vektor v_ℓ in $\mathcal{L}^{\text{hom}}$ enthalten ist, also $\phi_A(v_\ell) = Av_\ell = 0_{K^m}$ gilt. Daraus ergibt sich zumindest die Inklusion „ $\supseteq$ “ der angegebenen Gleichung, denn für alle $\lambda_\ell \in K$ mit $\ell \in S$ folgt dann mit Lemma (6.10) jeweils

$$A\left(\sum_{\ell \in S} \lambda_\ell v_\ell\right) = \phi_A\left(\sum_{\ell \in S} \lambda_\ell v_\ell\right) = \sum_{\ell \in S} \phi_A(\lambda_\ell v_\ell) = \sum_{\ell \in S} \lambda_\ell \phi_A(v_\ell) = \sum_{\ell \in S} \lambda_\ell \cdot 0_{K^m} = 0_{K^m}.$$

Später werden wir dann sehen, wie man anhand der ℓ -ten Komponente erkennt, dass die Vektoren v_ℓ linear unabhängig sind und somit einen $(n-r)$ -dimensionalen Untervektorraum des K^n bilden. Außerdem werden wir aus dem sog. *Dimensionssatz* für lineare Abbildungen herleiten, dass $\mathcal{L}^{\text{hom}}$ ebenfalls $(n-r)$ -dimensional ist, wodurch aus der Inklusion „ $\supseteq$ “ die Übereinstimmung der beiden Mengen folgt.

Seien nun $\ell \in S$ und $i \in \{1, \dots, m\}$ beliebig vorgegeben. Nach Definition der normierten Zeilenstufenform gilt $a_{ij_k} = \delta_{ik}$ für $1 \leq k \leq r$. Nach Definition des Vektors v_ℓ gilt $(v_\ell)_{j_k} = -a_{k\ell}$ für $1 \leq k \leq r$, $(v_\ell)_\ell = 1_K$ und $(v_\ell)_j = 0_K$ für alle übrigen $j \in \{1, \dots, r\}$. Es folgt

$$\begin{aligned} (Av_\ell)_i &= \sum_{j=1}^n a_{ij}(v_\ell)_j = a_{i\ell}(v_\ell)_\ell + \sum_{k=1}^r a_{ij_k}(v_\ell)_{j_k} = a_{i\ell} + \sum_{k=1}^r \delta_{ik} \cdot (v_\ell)_{j_k} \\ &= a_{i\ell} + (v_\ell)_{j_i} = a_{i\ell} + (-a_{i\ell}) = 0_K. \end{aligned}$$

Insgesamt gilt also tatsächlich $Av_\ell = 0_{K^m}$. □

Wir diskutieren eine Reihe von Anwendungsbeispielen für die soeben bewiesene Lösungsformel.

(i) Das homogene lineare Gleichungssystem $x_1 = 0$, $x_2 + 2x_3 = 0$ hat die Koeffizientenmatrix

$$A = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 2 \end{pmatrix}.$$

Es handelt sich um eine Matrix in normierter Zeilenstufenform mit den Kennzahlen $r = 2$, $j_1 = 1$, $j_2 = 2$. Es ist $S = \{1, 2, 3\} \setminus \{j_1, j_2\} = \{3\}$. Die Lösungsmenge ist somit $\mathcal{L}^{\text{hom}} = \{\lambda_3 v_3 \mid \lambda_3 \in \mathbb{R}\}$ mit dem Lösungsvektor

$$v_3 = \begin{pmatrix} 0 \\ -2 \\ 1 \end{pmatrix}.$$

(ii) Das homogene lineare Gleichungssystem $x_1 = 0$, $x_3 = 0$ hat die Koeffizientenmatrix

$$A = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Die Kennzahlen dieser normierten Zeilenstufenform lauten $r = 2$, $j_1 = 1$, $j_2 = 3$. Es ist $S = \{2\}$, und die Lösungsmenge ist gegeben durch $\mathcal{L}^{\text{hom}} = \{\lambda_2 v_2 \mid \lambda_2 \in \mathbb{R}\}$ mit

$$v_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}.$$

(iii) Das homogene lineare Gleichungssystem $x_1 - 4x_3 + 5x_5 = 0$, $x_2 + 2x_3 = 0$, $x_4 = 0$ hat die Koeffizientenmatrix

$$A = \begin{pmatrix} 1 & 0 & -4 & 0 & 5 \\ 0 & 1 & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \end{pmatrix},$$

eine normierte ZSF mit den Kennzahlen $r = 3$, $j_1 = 1$, $j_2 = 2$, $j_3 = 4$. Hier ist $S = \{1, \dots, 5\} \setminus \{j_1, j_2, j_3\} = \{3, 5\}$. Der Lösungsraum $\mathcal{L}^{\text{hom}} = \{\lambda_3 v_3 + \lambda_5 v_5 \mid \lambda_3, \lambda_5 \in K\}$ wird diesmal aufgespannt von den Vektoren

$$v_3 = \begin{pmatrix} 4 \\ -2 \\ 1 \\ 0 \\ 0 \end{pmatrix}, \quad v_5 = \begin{pmatrix} -5 \\ 0 \\ 0 \\ 0 \\ 1 \end{pmatrix}.$$

(7.5) Satz Sei $\tilde{A} = (A \ b) \in \mathcal{M}_{m \times (n+1), K}$ die erweiterte Koeffizientenmatrix eines LGS und $\mathcal{L} \subseteq K^n$ dessen Lösungsmenge. Wir setzen voraus, dass $\tilde{A}$ in normierter ZSF vorliegt, mit den Kennzahlen r und $j_1, \dots, j_r$.

(i) Ist $j_r = n + 1$, dann gilt $\mathcal{L} = \emptyset$.

(ii) Sei nun $j_r \leq n$. Wir definieren einen Vektor $w \in K^n$ durch $w = \sum_{k=1}^r b_k e_{j_k}$. Dann gilt $w \in \mathcal{L}$.

Der spezielle Lösungsvektor w entsteht also einfach dadurch, dass man die Werte $b_1, \dots, b_r$ auf die Positionen $j_1, \dots, j_r$ verteilt und die übrigen Komponenten auf Null setzt. Es gilt also $w_{j_k} = b_k$ für $1 \leq k \leq r$ und $w_\ell = 0$ für alle $\ell \in S$.

Beweis: zu (i) Nehmen wir an, dass $\mathcal{L}$ nichtleer und w ein Element aus $\mathcal{L}$ ist. Dann gilt insbesondere $\sum_{j=1}^n a_{rj} w_j = b_r$. Wegen $j_r = n + 1$ gilt aber $a_{rj} = 0_K$ für $1 \leq j \leq n$ und $b_r = a_{r, n+1} = a_{r, j_r} = 1_K$. Setzen wir dies in die Gleichung ein, so erhalten wir $\sum_{j=1}^n 0_K w_j = 1_K$. Der Widerspruch $0_K = 1_K$ zeigt, dass unsere Annahme falsch war.

zu (ii) Zu zeigen ist $\sum_{j=1}^n a_{kj} w_j = b_k$ für $1 \leq k \leq m$. Sei $k \in \{1, \dots, r\}$. Nach Definition der normierten ZSF gilt $a_{kj} = 0$ für $j < j_k$, und für $k \leq \ell \leq r$ ist $a_{\ell j_\ell} = 1$ der einzige Eintrag ungleich Null in der j_ℓ -ten Spalte. Es gilt also auch $a_{kj_\ell} = 0$ für $j > j_\ell$. Nach Definition des Vektors w erhalten wir für $1 \leq k \leq r$ somit

$$\sum_{j=1}^n a_{kj} w_j = \sum_{j=j_k}^n a_{kj} w_j = \sum_{\ell=k}^r a_{kj_\ell} b_\ell = a_{kj_k} b_k = 1 \cdot b_k = b_k.$$

Für $r < k \leq m$ gilt nach Eigenschaft (1) der normierten ZSF (Einträge unterhalb der r -ten Zeile gleich Null) sowohl $a_{kj} = 0$ für $1 \leq j \leq n$ als auch $b_k = 0$, also ebenfalls $\sum_{j=1}^n a_{kj} w_j = 0 = b_k$. Insgesamt ist die Gleichung $\sum_{j=1}^n a_{kj} w_j = b_k$ also tatsächlich für $1 \leq k \leq m$ erfüllt. $\square$

Wir demonstrieren die Anwendung der Lösungsformel an einem konkreten Beispiel. Das inhomogene LGS $x_1 - 4x_3 + 5x_5 = -2$, $x_2 + 2x_3 = 7$, $x_4 = 5$ besitzt die erweiterte Koeffizientenmatrix

$$\tilde{A} = \begin{pmatrix} 1 & 0 & -4 & 0 & 5 & -2 \\ 0 & 1 & 2 & 0 & 0 & 7 \\ 0 & 0 & 0 & 1 & 0 & 5 \end{pmatrix}.$$

Es handelt sich um eine normierte Zeilenstufenform mit den Kennzahlen $r = 3$ und $j_1 = 1, j_2 = 2, j_3 = 4$. Nach Satz (7.5) ist $w = (-2, 7, 0, 5, 0) \in \mathbb{R}^5$ ein Lösungsvektor, was man durch Einsetzen in die Gleichungen des Systems unmittelbar überprüft: Es ist $(-2) - 4 \cdot 0 + 5 \cdot 0 = -2$, $7 + 2 \cdot 0 = 7$ und $5 = 5$.

Die Lösung linearer Gleichungssysteme mit Koeffizientenmatrix in normierter Zeilenstufenform haben wir damit vollständig geklärt. Dass die Lösung linearer Gleichungssysteme mit *beliebiger* Koeffizientenmatrix darauf zurückgeführt werden kann, beruht auf der folgenden einfachen Beobachtung.

(7.6) Proposition Sei K ein Körper, und seien $m, n \in \mathbb{N}$. Sei $A \in \mathcal{M}_{m \times n, K}$ und $b \in K^m$. Dann gilt $\mathcal{L}_{A,b} = \mathcal{L}_{TA,Tb}$ für jede Matrix $T \in \text{GL}_m(K)$. Mit anderen Worten, die Lösungsmenge eines LGS ändert sich nicht, wenn man beide Seiten der Gleichung $Ax = b$ von links mit einer invertierbaren Matrix multipliziert.

Beweis: Für jeden Vektor $c \in K^n$ gilt die Äquivalenz

$$c \in \mathcal{L}_{A,b} \iff Ac = b \iff TAc = Tb \iff c \in \mathcal{L}_{TA,Tb}.$$

Dabei kommt die Richtung „ $\Leftarrow$ “ im zweiten Schritt durch die Rechnung

$$TAc = Tb \Rightarrow T^{-1}TAc = T^{-1}Tb \Rightarrow E^{(m)}Ac = E^{(m)}b \Rightarrow Ac = b$$

zu Stande. □

Man beachte, dass sich die Lösungsmenge bei Multiplikation der Gleichung $Ax = b$ mit einer *nicht-invertierbaren* Matrix durchaus verändern kann. Multipliziert man beide Seiten zum Beispiel mit der Nullmatrix $0^{(m)}$, dann erhält man die Gleichung $0^{(m \times n)}x = 0_{K^m}$. Die Lösungsmenge dieses linearen Gleichungssystems ist der gesamte K^n (weil jeder Vektor $c \in K^n$ die Gleichung $0^{(m \times n)}c = 0_{K^m}$ erfüllt), unabhängig davon, wie die Lösungsmenge von $Ax = b$ ausgesehen hat.

Um nun ein Lösungsverfahren für beliebige LGS zu erhalten, brauchen wir also nur noch ein Verfahren, mit dem wir beliebige Matrizen in normierte Zeilenstufenform überführen können. Dazu verwenden wir die Rechenoperationen für Matrizen, die in § 7 eingeführt wurden. Bei Rechnungen mit Matrizen ist es oft günstig, diese in mehrere Bereiche aufzuteilen. Sei $A \in \mathcal{M}_{m \times n, K}$, seien k_1, k_2, ℓ_1, ℓ_2 natürliche Zahlen mit $1 \leq k_1 \leq k_2 \leq m$, $1 \leq \ell_1 \leq \ell_2 \leq n$, und außerdem $r = k_2 - k_1 + 1$, $s = \ell_2 - \ell_1 + 1$. Dann nennt man die Matrix $B \in \mathcal{M}_{r \times s, K}$ mit den Einträgen $b_{ij} = a_{k_1+i-1, \ell_1+j-1}$ für $1 \leq i \leq r$, $1 \leq j \leq s$ eine **Teilmatrix** von A ; es handelt sich um einen „rechteckigen Ausschnitt“ der Matrix A .

Häufig verwendet man die sogenannte **Blockschreibweise**, um Matrizen darzustellen, die aus bestimmten Teilmatrizen aufgebaut sind. So steht beispielsweise der Ausdruck

$$\begin{pmatrix} A & B \\ C & D \end{pmatrix}$$

für die Matrix, deren linker oberer Teil aus den Einträgen von A und entsprechend in den übrigen drei Bereichen aus den Einträgen von B , C und D besteht. Dabei wird vorausgesetzt, dass untereinander stehende Matrizen (hier: A, C bzw. B, D) stets dieselbe Spaltenzahl und nebeneinander stehende Matrizen (A, B bzw. C, D) dieselbe Zeilenzahl haben. Das Rechnen mit Matrizen in Blockschreibweise wird durch eine Reihe von Rechenregeln vereinfacht.

(7.7) Proposition Seien A, B, C, D Matrizen über K .

(i) Stimmt die Spaltenzahl von A und B mit der Zeilenzahl von C überein, dann gilt

$$\begin{pmatrix} A \\ B \end{pmatrix} C = \begin{pmatrix} AC \\ BC \end{pmatrix}.$$

(ii) Stimmt die Spaltenzahl von A mit der Zeilenzahl von B und C überein, dann gilt

$$A \begin{pmatrix} B & C \end{pmatrix} = \begin{pmatrix} AB & AC \end{pmatrix}.$$

(iii) Stimmt die Spaltenzahl von A mit der Zeilenzahl von C und die Spaltenzahl von B mit der Zeilenzahl von D überein, dann gilt

$$\begin{pmatrix} A & B \end{pmatrix} \begin{pmatrix} C \\ D \end{pmatrix} = AC + BD.$$

Beweis: Wir beschränken uns auf den Beweis von (iii). Nach Voraussetzung gilt $A \in \mathcal{M}_{m \times n_1, K}$, $B \in \mathcal{M}_{m \times n_2, K}$, $C \in \mathcal{M}_{n_1 \times r, K}$ und $D \in \mathcal{M}_{n_2 \times r, K}$ für geeignete $m, n_1, n_2, r \in \mathbb{N}$. Die Matrix $AC + BD$ auf der rechten Seite ist in $\mathcal{M}_{m \times r, K}$ enthalten. Seien nun k, ℓ mit $1 \leq k \leq m$ und $1 \leq \ell \leq r$ vorgegeben. Zu zeigen ist, dass der Eintrag des Matrixprodukts links an der Position (k, ℓ) mit dem Eintrag der Matrix $AC + BD$ an derselben Stelle übereinstimmt. Um den Eintrag auf der linken Seite auszurechnen, muss die k -te Zeile des Faktors $\begin{pmatrix} A & B \end{pmatrix}$ mit der ℓ -ten Spalte des zweiten Faktors $\begin{pmatrix} C \\ D \end{pmatrix}$ multipliziert werden. Dies liefert den Wert

$$\sum_{j=1}^{n_1} a_{kj} c_{j\ell} + \sum_{j=1}^{n_2} b_{kj} d_{j\ell}.$$

Die erste Summe entspricht dem Eintrag von AC an der Stelle (k, ℓ) , die zweite Summe dem Eintrag von BD an derselben Position. Insgesamt erhalten wir also den Eintrag von $AC + BD$ an der Stelle (k, ℓ) . $\square$

Wir demonstrieren die Funktionsweise der Rechenregel (7.7) (iii) für Blockmatrizen anhand eines Beispiels und betrachten dazu die vier Matrizen

$$A = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}, \quad B = \begin{pmatrix} 5 & 6 \\ 7 & 8 \end{pmatrix}, \quad C = \begin{pmatrix} 9 & 10 \\ 11 & 12 \end{pmatrix}, \quad D = \begin{pmatrix} 13 & 14 \\ 15 & 16 \end{pmatrix}.$$

Es gilt

$$AC = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \begin{pmatrix} 9 & 10 \\ 11 & 12 \end{pmatrix} = \begin{pmatrix} 31 & 34 \\ 71 & 78 \end{pmatrix} \quad \text{und} \quad BD = \begin{pmatrix} 5 & 6 \\ 7 & 8 \end{pmatrix} \begin{pmatrix} 13 & 14 \\ 15 & 16 \end{pmatrix} = \begin{pmatrix} 155 & 166 \\ 211 & 226 \end{pmatrix}$$

und somit

$$AC + BD = \begin{pmatrix} 155 & 166 \\ 211 & 226 \end{pmatrix} + \begin{pmatrix} 31 & 34 \\ 71 & 78 \end{pmatrix} = \begin{pmatrix} 186 & 200 \\ 282 & 304 \end{pmatrix}.$$

Eine direkte Multiplikation der zusammengesetzten Matrizen liefert dasselbe Ergebnis:

$$\begin{pmatrix} AB \end{pmatrix} \begin{pmatrix} C \\ D \end{pmatrix} = \begin{pmatrix} 1 & 2 & 5 & 6 \\ 3 & 4 & 7 & 8 \end{pmatrix} \begin{pmatrix} 9 & 10 \\ 11 & 12 \\ 13 & 14 \\ 15 & 16 \end{pmatrix} = \begin{pmatrix} 186 & 200 \\ 282 & 304 \end{pmatrix}.$$

Allgemeiner kann gezeigt werden, dass man Matrizen mit beliebiger Aufteilung „blockweise“ multiplizieren kann, wobei lediglich vorausgesetzt werden muss, dass die Teilmatrizen, die dabei multipliziert werden sollen, „zusammenpassen“.

(7.8) Proposition Seien $m, n, r \in \mathbb{N}$, seien $A^{(i,j)}$ für $1 \leq i \leq m$, $1 \leq j \leq n$ und $B^{(j,k)}$ für $1 \leq j \leq n$ und $1 \leq k \leq r$ Matrizen mit der Eigenschaft, dass die Spaltenzahl von $A^{(i,j)}$ jeweils mit der Zeilenzahl von $B^{(j,k)}$ übereinstimmt, für alle i, j, k . Außerdem setzen wir voraus, dass die Zeilenzahlen von $A^{(i,j)}$ für festes i und die Spaltenzahlen von $B^{(j,k)}$ für festes k jeweils gleich sind. Dann gilt

$$\begin{pmatrix} A^{(1,1)} & \dots & A^{(1,n)} \\ \vdots & & \vdots \\ A^{(m,1)} & \dots & A^{(m,n)} \end{pmatrix} \begin{pmatrix} B^{(1,1)} & \dots & B^{(1,r)} \\ \vdots & & \vdots \\ B^{(n,1)} & \dots & B^{(n,r)} \end{pmatrix} = \begin{pmatrix} C^{(1,1)} & \dots & C^{(1,r)} \\ \vdots & & \vdots \\ C^{(m,1)} & \dots & C^{(m,r)} \end{pmatrix}$$

mit $C^{(i,k)} = \sum_{j=1}^n A^{(i,j)} B^{(j,k)}$ für $1 \leq i \leq m$ und $1 \leq k \leq r$.

Beweis: Wir geben den Beweis nur der Vollständigkeit halber an, für den weiteren Verlauf ist er ohne Belang. Für alle i, j, k sei $m_i \times n_j$ jeweils das Format der Matrix $A^{(i,j)}$ und $n_j \times r_k$ das Format der Matrix $B^{(j,k)}$. Dann hat die Matrix $C^{(i,k)}$ jeweils das Format $m_i \times r_k$. Wir bezeichnen die Matrix auf der rechten Seite der zu beweisenden Gleichung mit D und die Matrix auf der linken Seite mit C . Beide Matrizen haben das Format $m_0 \times r_0$ mit $m_0 = \sum_{i=1}^m m_i$ und $r_0 = \sum_{k=1}^r r_k$. Außerdem sei A die Matrix mit den Blöcken $A^{(i,j)}$ und B die Matrix mit den Blöcken $B^{(j,k)}$. Nach Definition gilt $D = AB$.

Seien nun $p \in \{1, \dots, m_0\}$ und $q \in \{1, \dots, r_0\}$ vorgegeben. Zu zeigen ist $c_{pq} = d_{pq}$. Der Eintrag d_{pq} kommt dadurch zu Stande, dass die p -te Zeile von A mit der q -ten Spalte von B multipliziert wird. Wir nehmen nun an, dass $i \in \{1, \dots, m\}$ und $k \in \{1, \dots, r\}$ so gewählt sind, dass die p -te Zeile der Matrix A durch die f -ten Zeilen der Matrizen $A^{(i,1)}, A^{(i,2)}, \dots, A^{(i,n)}$ läuft, und ebenso die q -te Spalte von B durch die g -ten Spalten der Matrizen $B^{(1,k)}, B^{(2,k)}, \dots, B^{(n,k)}$. Dabei ist $f \in \{1, \dots, m_i\}$ und $g \in \{1, \dots, r_k\}$. Setzen wir $n_0 = \sum_{j=1}^n n_j$, dann gilt

$$d_{pq} = \sum_{j=1}^{n_0} a_{pj} b_{jq} = \sum_{j=1}^n \sum_{\ell=1}^{n_j} a_{f\ell}^{(i,j)} b_{\ell g}^{(j,k)}.$$

Nun läuft die p -te Zeile von C auch durch die f -ten Zeilen der Matrizen $C^{(i,1)}, C^{(i,2)}, \dots, C^{(i,n)}$, und die q -te Spalte von C entsprechend durch die g -ten Spalten der Matrizen $C^{(1,k)}, C^{(2,k)}, \dots, C^{(n,k)}$. Wegen $C^{(i,k)} = \sum_{j=1}^n A^{(i,j)} B^{(j,k)}$ für $1 \leq i \leq m$ und $1 \leq k \leq r$ erhalten wir dann wie gewünscht

$$c_{pq} = c_{fg}^{(i,k)} = \sum_{j=1}^n (A^{(i,j)} B^{(j,k)})_{fg} = \sum_{j=1}^n \sum_{\ell=1}^{n_j} a_{f\ell}^{(i,j)} b_{\ell g}^{(j,k)} = d_{pq}. \quad \square$$

(7.9) Definition Eine Matrix aus $\mathcal{M}_{m,K}$ der Form $M_{k,\lambda} = E^{(m)} + (\lambda - 1)B_{kk}^{(m \times m)}$ mit $k \in \{1, \dots, m\}$ und $\lambda \in K^\times$ oder der Form $A_{k,\ell,\lambda} = E^{(m)} + \lambda B_{\ell k}^{(m \times m)}$ mit $k, \ell \in \{1, \dots, m\}$ und $\lambda \in K$ wird **Elementarmatrix** genannt.

In Blockschreibweise hat die Elementarmatrix $M_{k,\lambda}$ die Form

$$M_{k,\lambda} = \begin{pmatrix} E^{(k-1)} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \lambda & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & E^{(m-k)} \end{pmatrix}$$

wobei die Einträge $\mathbf{0}$ jeweils für Nullmatrizen der passenden Größe stehen. Die Elementarmatrix $A_{k,\ell,\lambda}$ hat im Fall $k < \ell$ bzw. $k > \ell$ die Form

$$A_{k,\ell,\lambda} = \begin{pmatrix} E^{(k-1)} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & 1 & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & E^{(\ell-k-1)} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \lambda & \mathbf{0} & 1 & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & E^{(m-\ell)} \end{pmatrix}$$

beziehungsweise

$$A_{k,\ell,\lambda} = \begin{pmatrix} E^{(\ell-1)} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & 1 & \mathbf{0} & \lambda & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & E^{(k-\ell-1)} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & 1 & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & E^{(m-k)} \end{pmatrix}$$

(7.10) Proposition Sei $A \in \mathcal{M}_{m \times n, K}$.

- (i) Sei $\lambda \in K^\times$ und $k \in \{1, \dots, m\}$. Multipliziert man die Matrix A von links mit der Elementarmatrix $M_{k,\lambda}$, so bewirkt dies eine Multiplikation der k -ten Zeile mit dem Wert λ .
- (ii) Seien $k, \ell \in \{1, \dots, m\}$ mit $k \neq \ell$ und $\lambda \in K$. Multipliziert man die Matrix A mit der Elementarmatrix $A_{k,\ell,\lambda}$, dann wird das λ -fache der k -ten Zeile zur ℓ -ten Zeile von A addiert.

Beweis: zu (i) Sei $B \in \mathcal{M}_{(k-1) \times n, K}$ die Teilmatrix bestehend aus den oberen $k-1$ und $C \in \mathcal{M}_{(m-k) \times n, K}$ die Teilmatrix bestehend aus den unteren $m-k$ Zeilen von A . Ferner sei $z \in \mathcal{M}_{1 \times n, K}$ die k -te Zeile von A . Dann gilt

$$M_{k, \lambda} A = \begin{pmatrix} E^{(k-1)} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \lambda & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & E^{(m-k)} \end{pmatrix} \begin{pmatrix} B \\ z \\ C \end{pmatrix} = \begin{pmatrix} E^{(k-1)}B + \mathbf{0}z + \mathbf{0}C \\ \mathbf{0}B + \lambda z + \mathbf{0}C \\ \mathbf{0}B + \mathbf{0}z + E^{(m-k)}C \end{pmatrix} = \begin{pmatrix} B \\ \lambda z \\ C \end{pmatrix}$$

zu (ii) Hier beschränken wir uns auf den Fall $k < \ell$ und teilen die Matrix A auf in die Matrix $B \in \mathcal{M}_{(k-1) \times n, K}$ bestehend aus den ersten $k-1$ Zeilen, der Matrix $C \in \mathcal{M}_{(\ell-k-1) \times n, K}$ bestehend aus der $(k+1)$ -ten bis zur $(\ell-1)$ -ten Zeile und der Matrix $D \in \mathcal{M}_{(m-\ell) \times n, K}$ bestehend aus den unteren $m-\ell$ Zeilen. Ferner seien $z_k, z_\ell \in \mathcal{M}_{1 \times n, K}$ die k -te und ℓ -te Zeile von A . Dann erhalten wir

$$A_{k, \ell, \lambda} A = \begin{pmatrix} E^{(k-1)} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & 1 & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & E^{(\ell-k-1)} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \lambda & \mathbf{0} & 1 & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & E^{(m-\ell)} \end{pmatrix} \begin{pmatrix} B \\ z_k \\ C \\ z_\ell \\ D \end{pmatrix} = \begin{pmatrix} B \\ z_k \\ C \\ \lambda z_k + z_\ell \\ D \end{pmatrix}$$

□

Wir zeigen anhand zweier Beispiele, dass die Multiplikation mit Elementarmatrizen tatsächlich den angegebenen Effekt hat. Die Multiplikation einer dreizeiligen Matrix mit $M_{2,3}$ von links bewirkt eine Multiplikation der zweiten Zeile mit dem Wert 3. Zum Beispiel gilt

$$M_{2,3} \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 3 & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 15 & 18 & 21 & 24 \\ 9 & 10 & 11 & 12 \end{pmatrix}.$$

Ebenso bewirkt die Multiplikation mit $A_{1,3,2}$ von links, dass das zweifache der ersten Zeile zur dritten addiert wird, zum Beispiel

$$A_{1,3,2} \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 2 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 9 & 10 & 11 & 12 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 5 & 6 & 7 & 8 \\ 11 & 14 & 17 & 20 \end{pmatrix}.$$

Die in der Proposition beschriebenen Umformungen werden **elementare Zeilenumformungen** genannt. Jede elementare Zeilenumformung einer Matrix A lässt sich also durch Multiplikation mit einer Elementarmatrix von links realisieren. Dementsprechend führt die Multiplikation von A mit einem Produkt $E_m \cdot E_{m-1} \cdot \dots \cdot E_1$ von Elementarmatrizen dazu, dass A einer Folge von m Zeilenumformungen unterworfen wird. Wir bezeichnen die Menge aller Matrizen in $\mathcal{M}_{m, K}$, die sich als Produkt von Elementarmatrizen schreiben lassen, mit $\mathcal{E}_m(K)$. Übrigens lässt sich auch die Vertauschung von Zeilen durch eine Folge von elementaren Umformungen wie oben beschrieben bewerkstelligen, wie man an dem Schema

$$\begin{pmatrix} a_{k\bullet} \\ a_{\ell\bullet} \end{pmatrix} \xrightarrow{A_{k, \ell, 1}} \begin{pmatrix} a_{k\bullet} \\ a_{k\bullet} + a_{\ell\bullet} \end{pmatrix} \xrightarrow{M_{k, -1}} \begin{pmatrix} -a_{k\bullet} \\ a_{k\bullet} + a_{\ell\bullet} \end{pmatrix} \xrightarrow{A_{\ell, k, 1}} \begin{pmatrix} a_{\ell\bullet} \\ a_{k\bullet} + a_{\ell\bullet} \end{pmatrix} \xrightarrow{A_{k, \ell, -1}} \begin{pmatrix} a_{\ell\bullet} \\ a_{k\bullet} \end{pmatrix}$$

erkennt. Mit Hilfe des Matrixkalküls werden wir nun zeigen, dass sich jede Matrix durch eine endliche Anzahl von Zeilenumformungen auf normierte Zeilenstufenform bringen lässt.

(7.11) Lemma Sei $A \in \mathcal{M}_{m \times 1, K}$ eine Matrix, die aus einer einzigen Spalte besteht, also eine Matrix der Form $A = {}^t(a_1 \ a_2 \ \dots \ a_m)$. Sind nicht alle Einträge von A gleich Null, dann gibt es ein Produkt $E \in \mathcal{E}_m(K)$ von Elementarmatrizen mit $EA = {}^t(1_K \ 0_K \ 0_K \ \dots \ 0_K)$.

Beweis: Auf Grund unserer Vorbemerkung genügt es zu zeigen, dass A durch eine endliche Abfolge von elementaren Zeilenumformungen auf die Gestalt ${}^t(1 \ 0 \ \dots \ 0)$ gebracht werden kann. Auch Vertauschungen von Zeilen sind zulässig, weil diese (wie oben gesehen) durch endlich viele elementare Umformungen realisierbar sind. Nach Voraussetzung gibt es ein $k \in \{1, \dots, m\}$ mit $a_k \neq 0_K$. Nach Multiplikation der k -ten Zeile mit a_k^{-1} und Vertauschung der k -ten mit der ersten Zeile gilt $a_1 = 1_K$. Nun addieren wir für $\ell = 2, \dots, m$ jeweils das $(-a_\ell)$ -fache der ersten Zeile zur ℓ -ten. Dies führt dazu, dass sämtliche Einträge der Matrix mit Ausnahme des ersten zu Null werden. $\square$

(7.12) Satz Jede Matrix $A \in \mathcal{M}_{m \times n, K}$ kann durch endlich viele elementare Zeilenumformungen auf normierte ZSF gebracht werden. Eine äquivalente Formulierung dieser Aussage lautet: Es gibt eine Matrix $E \in \mathcal{E}_m(K)$, so dass EA in normierter ZSF vorliegt.

Beweis: Wir zeigen zunächst, dass A auf ZSF gebracht werden kann und führen den Beweis durch vollständige Induktion über die Anzahl n der Spalten. Der Fall $n = 1$ ist mit Lemma (7.11) bereits erledigt, denn nach Definition ist ${}^t(1_K \ 0_K \ \dots \ 0_K)$ eine Matrix in ZSF (mit den Kennzahlen $r = j_1 = 1$). Sei nun $n \in \mathbb{N}$, und setzen wir die Aussage für dieses n voraus. Sei außerdem $A \in \mathcal{M}_{m \times (n+1), K}$ eine beliebige Matrix. Wir müssen zeigen, dass A auf ZSF gebracht werden kann und unterscheiden dafür zwei Fälle.

1. Fall: Die erste Spalte von A hat nur Nulleinträge.

Dann hat A die Form $(\mathbf{0}^{(m \times 1)} \ B)$ mit einer Matrix $B \in \mathcal{M}_{m \times n, K}$. Nach Induktionsvoraussetzung gibt es eine Matrix $E \in \mathcal{E}_m(K)$, so dass $B' = EB$ in ZSF vorliegt, mit gewissen Kennzahlen $r, j_1, \dots, j_r$. Es gilt

$$EA = E(\mathbf{0}^{(m \times 1)} \ B) = (\mathbf{0}^{(m \times 1)} \ EB) = (\mathbf{0}^{(m \times 1)} \ B').$$

Wie man leicht überprüft, liegt auch $(\mathbf{0}^{(m \times 1)} \ B')$ die Matrix in ZSF vor, mit den Kennzahlen $r, j_1 + 1, \dots, j_r + 1$.

2. Fall: Die erste Spalte von A hat Einträge ungleich Null.

In diesem Fall kann A in der Blockgestalt

$$A = \begin{pmatrix} a_{11} & z \\ s & C \end{pmatrix}$$

dargestellt werden, mit $a_{11} \in K$, $z \in \mathcal{M}_{1 \times n, K}$, $s \in \mathcal{M}_{(m-1) \times 1, K}$ und $C \in \mathcal{M}_{(m-1) \times n, K}$, wobei die Teilmatrix ${}^t(a_{11} \ s)$ nicht nur Nulleinträge enthält. Nach Lemma (7.11) gibt es eine Matrix $E \in \mathcal{E}_m(K)$ mit

$$E \begin{pmatrix} a_{11} \\ s \end{pmatrix} = \begin{pmatrix} 1 \\ \mathbf{0} \end{pmatrix}$$

und wir erhalten

$$EA = \begin{pmatrix} 1 & z' \\ \mathbf{0} & C' \end{pmatrix}$$

mit geeigneten Matrizen $z' \in \mathcal{M}_{1 \times n, K}$ und $C' \in \mathcal{M}_{(m-1) \times n, K}$. Nach Induktionsvoraussetzung existiert nun eine Matrix $E' \in \mathcal{E}_{m-1}(K)$, so dass $E'C'$ in ZSF vorliegt, mit gewissen Kennzahlen $r, j_1, \dots, j_r$. Außerdem gilt

$$\begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & E' \end{pmatrix} \begin{pmatrix} 1 & z' \\ \mathbf{0} & C' \end{pmatrix} = \begin{pmatrix} 1 & z' \\ \mathbf{0} & E'C' \end{pmatrix}$$

Wieder überprüft man, dass sich die Matrix rechts in ZSF befindet, mit Kennzahlen $r+1, 1, j_1+1, \dots, j_r+1$. Anhand der Gleichung

$$\begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & U \end{pmatrix} \begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & V \end{pmatrix} = \begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & UV \end{pmatrix}$$

für Blockmatrizen sieht man, dass mit E' auch die Matrix

$$\begin{pmatrix} 1 & \mathbf{0} \\ \mathbf{0} & E' \end{pmatrix}$$

als Produkt von Elementarmatrizen darstellbar ist.

Zu zeigen bleibt, dass jede Matrix in ZSF durch elementare Zeilenumformungen auf *normierte* ZSF gebracht werden kann. Dazu setzen wir voraus, dass A bereits in ZSF mit den Kennzahlen $r, j_1, \dots, j_r$ vorliegt. Um $a_{ij_i} = 1$ für $1 \leq i \leq r$ zu erreichen, dividiert man einfach für jedes i die i -te Zeile durch a_{ij_i} . Die ZSF der Matrix wird durch diese Operation nicht zerstört, da die Eigenschaft eines Eintrages, gleich Null oder ungleich Null zu sein, dadurch nicht verändert wird.

Die Bedingung $a_{kj_i} = 0$ für $k < i$ kann dadurch erfüllt werden, dass man nacheinander für die Zeilennummern $i = r, r-1, r-2, \dots, 1$ jeweils das a_{kj_i} -fache der i -ten Zeile von der k -ten Zeile subtrahiert, für $1 \leq k < i$. Dabei ist darauf zu achten, dass in keinem Schritt die ZSF beeinträchtigt wird und die erreichte Form für die Spalten j_ℓ mit $\ell > i$ erhalten bleibt. Die ZSF bleibt erhalten, da die i -te Zeile ihren ersten Eintrag $\neq 0$ erst in der Spalte j_i hat und $j_i > j_k$ für $1 \leq k < i$ gilt. Somit werden weder die Zeilenköpfe der darüberliegenden Zeilen noch die Einträge links davon verändert. Die Zeilen unterhalb der i -ten bleiben völlig unverändert. Auch die Bedingung $a_{kj_\ell} = 0$ für $\ell > i$ und $k < \ell$ bleibt erhalten, da der einzige Eintrag ungleich Null in der j_ℓ -ten Spalte der Eintrag $a_{\ell j_\ell} = 1$ ist, und dieser spielt wegen $\ell > i$ bei der Zeilenumformung keine Rolle. $\square$

Man kann sich an diesem „induktiven“ Beweisschema orientieren, um eine beliebige, konkret vorgegebene Matrix $A \in \mathcal{M}_{m \times n, K}$ zunächst auf Zeilenstufenform und dann auf normierte Zeilenstufenform zu bringen.

Damit haben wir nun insgesamt ein vollständiges Verfahren zur Verfügung, um die Lösungsmenge eines beliebigen linearen Gleichungssystems der Form $Ax = b$ (mit $A \in \mathcal{M}_{m \times n, K}$ und $b \in K^n$) vollständig zu bestimmen. Dieses Verfahren wird auch als **Gauß'sches Eliminationsverfahren** bezeichnet.

Zunächst bringt man die erweiterte Koeffizientenmatrix $\tilde{A} = \begin{pmatrix} A & b \end{pmatrix} \in \mathcal{M}_{m \times (n+1), K}$ durch eine Folge von Zeilenumformungen auf normierte Zeilenstufenform. Durch Proposition (7.6) ist gewährleistet, dass sich die Lösungsmenge des Systems durch die Umformungen nicht ändert. An der umgeformten Matrix liest man eine spezielle Lösung $c \in L_{A,b}$ des Systems ab, und an der linken $m \times n$ -Teilmatrix den Lösungsraum $\mathcal{L}_{A,b}^{\text{hom}}$. Wie oben erläutert, gilt dann $\mathcal{L}_{A,b} = c + \mathcal{L}_{A,b}^{\text{hom}}$.

Wir illustrieren die Vorgehensweise anhand der beiden linearen Gleichungssysteme

$$\begin{array}{rclcl} 3x & & + 2z & = & 8 & & 3x & + & 5y & - & 3z & = & -4 \\ -x & + & 2y & + & 5z & = & -3 & \text{und} & 2x & - & 8y & + & 7z & = & 11 \\ & & 7y & - & 2z & = & -23 & & 5x & - & 3y & + & 4z & = & 7 \end{array}$$

aus § 6. Die erweiterte Koeffizientenmatrix des ersten Systems wird durch die Schritte

$$\begin{aligned} \begin{pmatrix} 3 & 0 & 2 & 8 \\ -1 & 2 & 5 & -3 \\ 0 & 7 & -2 & -23 \end{pmatrix} &\rightarrow \begin{pmatrix} 1 & -2 & -5 & 3 \\ 3 & 0 & 2 & 8 \\ 0 & 7 & -2 & -23 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & -2 & -5 & 3 \\ 0 & 6 & 17 & -1 \\ 0 & 7 & -2 & -23 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & -2 & -5 & 3 \\ 0 & 6 & 17 & -1 \\ 0 & 1 & -19 & -22 \end{pmatrix} \rightarrow \\ \begin{pmatrix} 1 & -2 & -5 & 3 \\ 0 & 1 & -19 & -22 \\ 0 & 6 & 17 & -1 \end{pmatrix} &\rightarrow \begin{pmatrix} 1 & -2 & -5 & 3 \\ 0 & 1 & -19 & -22 \\ 0 & 0 & 131 & 131 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & -2 & -5 & 3 \\ 0 & 1 & -19 & -22 \\ 0 & 0 & 1 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & -2 & 0 & 8 \\ 0 & 1 & 0 & -3 \\ 0 & 0 & 1 & 1 \end{pmatrix} \\ &\rightarrow \begin{pmatrix} 1 & 0 & 0 & 2 \\ 0 & 1 & 0 & -3 \\ 0 & 0 & 1 & 1 \end{pmatrix} \end{aligned}$$

auf normierte Zeilenstufenform gebracht. An der umgeformten Matrix kann die spezielle Lösung $(2, -3, 1)$ abgelesen werden. Auch die linke 3×3 -Teilmatrix befindet sich in normierter Zeilenstufenform, mit den Kennzahlen $r = 3$, $j_1 = 1$, $j_2 = 2$, $j_3 = 3$. Mit der Notation von oben gilt $S = \emptyset$. Dies zeigt, dass $\mathcal{L}_{A,b}^{\text{hom}} = \{(0, 0, 0)\}$ gilt und somit die gesamte Lösungsmenge durch $\mathcal{L}_{A,b} = \{(2, -3, 1)\}$ gegeben ist.

Beim zweiten System ergeben die Umformungsschritte

$$\begin{aligned} \begin{pmatrix} 3 & 5 & -3 & -4 \\ 2 & -8 & 7 & 11 \\ 5 & -3 & 4 & 7 \end{pmatrix} &\rightarrow \begin{pmatrix} 1 & 13 & -10 & -15 \\ 2 & -8 & 7 & 11 \\ 5 & -3 & 4 & 7 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 13 & -10 & -15 \\ 0 & -34 & 27 & 41 \\ 0 & -68 & 54 & 82 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 13 & -10 & -15 \\ 0 & 1 & -\frac{27}{34} & -\frac{41}{34} \\ 0 & -68 & 54 & 82 \end{pmatrix} \\ &\rightarrow \begin{pmatrix} 1 & 13 & -10 & -15 \\ 0 & 1 & -\frac{27}{34} & -\frac{41}{34} \\ 0 & 0 & 0 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 13 & -10 & -15 \\ 0 & 1 & -\frac{27}{34} & -\frac{41}{34} \\ 0 & 0 & 0 & 0 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & \frac{11}{34} & \frac{23}{34} \\ 0 & 1 & -\frac{27}{34} & -\frac{41}{34} \\ 0 & 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

Hier lesen wir mit dem Verfahren von oben die spezielle Lösung $(\frac{23}{34}, -\frac{41}{34}, 0)$ ab. Die Kennzahlen der normierten ZSF in der linken 3×3 -Teilmatrix lauten $r = 2$, $j_1 = 1$, $j_2 = 2$. Es gilt also $S = \{3\}$, und an der Teilmatrix kann der Basisvektor $v_3 = (-\frac{11}{34}, \frac{27}{34}, 1)$ des Lösungsraums $\mathcal{L}_{A,b}^{\text{hom}}$ abgelesen werden. Insgesamt erhalten wir die Lösungsmenge

$$\mathcal{L}_{A,b} = \left\{ \begin{pmatrix} \frac{23}{34} \\ -\frac{41}{34} \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} -\frac{11}{34} \\ \frac{27}{34} \\ 1 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} 1 \\ -2 \\ -1 \end{pmatrix} + \lambda \begin{pmatrix} -\frac{11}{34} \\ \frac{27}{34} \\ 1 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} 1 \\ -2 \\ -1 \end{pmatrix} + \lambda \begin{pmatrix} -11 \\ 27 \\ 34 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\}$$

Hier haben die Darstellung der Lösungsmenge in den beiden folgenden Schritten etwas optimiert: Im ersten Schritt haben wir verwendet, dass neben $(\frac{23}{34}, -\frac{41}{34}, 0)$ auch $(1, -2, -1)$ eine spezielle Lösung des Systems ist, wie man durch Setzen von $\lambda = -1$ sieht. Offenbar ändert sich die Lösungsmenge auch nicht, wenn man den Basisvektor v_3 mit einem beliebigen skalaren Vielfachen multipliziert. In diesem Fall haben wir durch Multiplikation mit 34 die Bruchzahlen in v_3 beseitigt.

Wir beschäftigen uns nun noch mit der Frage, wie man die Invertierbarkeit von Matrizen nachweist, und wie gegebenenfalls die inverse Matrix berechnet werden kann.

(7.13) Satz Lässt sich eine Matrix $A \in \mathcal{M}_{n,K}$ durch endliche viele elementare Zeilenumformungen in eine Matrix A' in normierter ZSF mit Zeilenrang $r = n$ umwandeln, so ist A invertierbar.

Beweis: Bereits oben haben wir bemerkt, dass eine Matrix A' in normierter ZSF mit Zeilenrang $r = n$ zwangsläufig mit der Einheitsmatrix $E^{(n)}$ übereinstimmt. Weil A durch elementare Zeilenumformungen in $A' = E^{(n)}$ überführt werden kann, gibt es eine Matrix $T \in \mathcal{E}_n(K) \subseteq \text{GL}_n(K)$ mit $TA = E^{(n)}$. Es folgt $A = E^{(n)}A = (T^{-1}T)A = T^{-1}(TA) = T^{-1}E^{(n)} = T^{-1}$. Damit ist die Invertierbarkeit von A bewiesen. $\square$

Die Beweisidee in Satz (7.13) kann genutzt werden, um die zu A inverse Matrix auszurechnen. Wendet man die Zeilenumformungen im Beweis statt auf A auf die Blockmatrix $(A \ E^{(n)})$ an, so erhält man die Matrix

$$T(A \ E^{(n)}) = (TA \ TE^{(n)}) = (E^{(n)} \ T).$$

Aus der rechten Hälfte der umgeformten Matrix kann die Inverse von A abgelesen werden, denn es gilt die Äquivalenz $A = T^{-1} \Leftrightarrow A^{-1} = T$. Wir demonstrieren dieses Berechnungsverfahren, indem wir A^{-1} für die Matrix

$$A = \begin{pmatrix} 3 & 0 & 2 \\ -1 & 1 & -1 \\ 7 & 0 & 5 \end{pmatrix}$$

bestimmen. Dazu schreiben wir die Einheitsmatrix $E^{(3)}$ neben unsere Matrix A und formen auf normierte ZSF um.

$$\begin{aligned} & \left(\begin{array}{ccc|ccc} 3 & 0 & 2 & 1 & 0 & 0 \\ -1 & 1 & -1 & 0 & 1 & 0 \\ 7 & 0 & 5 & 0 & 0 & 1 \end{array} \right) \rightarrow \left(\begin{array}{ccc|ccc} -1 & 1 & -1 & 0 & 1 & 0 \\ 3 & 0 & 2 & 1 & 0 & 0 \\ 7 & 0 & 5 & 0 & 0 & 1 \end{array} \right) \rightarrow \\ & \left(\begin{array}{ccc|ccc} 1 & -1 & 1 & 0 & -1 & 0 \\ 3 & 0 & 2 & 1 & 0 & 0 \\ 7 & 0 & 5 & 0 & 0 & 1 \end{array} \right) \rightarrow \left(\begin{array}{ccc|ccc} 1 & -1 & 1 & 0 & -1 & 0 \\ 0 & 3 & -1 & 1 & 3 & 0 \\ 0 & 7 & -2 & 0 & 7 & 1 \end{array} \right) \rightarrow \\ & \left(\begin{array}{ccc|ccc} 1 & -1 & 1 & 0 & -1 & 0 \\ 0 & 1 & -\frac{1}{3} & \frac{1}{3} & 1 & 0 \\ 0 & 7 & -2 & 0 & 7 & 1 \end{array} \right) \rightarrow \left(\begin{array}{ccc|ccc} 1 & -1 & 1 & 0 & -1 & 0 \\ 0 & 1 & -\frac{1}{3} & \frac{1}{3} & 1 & 0 \\ 0 & 0 & \frac{1}{3} & -\frac{7}{3} & 0 & 1 \end{array} \right) \rightarrow \\ & \left(\begin{array}{ccc|ccc} 1 & -1 & 1 & 0 & -1 & 0 \\ 0 & 1 & -\frac{1}{3} & \frac{1}{3} & 1 & 0 \\ 0 & 0 & 1 & -7 & 0 & 3 \end{array} \right) \rightarrow \left(\begin{array}{ccc|ccc} 1 & -1 & 0 & 7 & -1 & -3 \\ 0 & 1 & 0 & -2 & 1 & 1 \\ 0 & 0 & 1 & -7 & 0 & 3 \end{array} \right) \rightarrow \left(\begin{array}{ccc|ccc} 1 & 0 & 0 & 5 & 0 & -2 \\ 0 & 1 & 0 & -2 & 1 & 1 \\ 0 & 0 & 1 & -7 & 0 & 3 \end{array} \right) \end{aligned}$$

Als Ergebnis erhalten wir also

$$A^{-1} = \begin{pmatrix} 3 & 0 & 2 \\ -1 & 1 & -1 \\ 7 & 0 & 5 \end{pmatrix}^{-1} = \begin{pmatrix} 5 & 0 & -2 \\ -2 & 1 & 1 \\ -7 & 0 & 3 \end{pmatrix}.$$

Offen bleibt hierbei die Frage, wie es zu interpretieren ist, wenn die Matrix A zwar auf normierte ZSF, aber mit Zeilenrang $r < n$, gebracht werden kann. Dafür ist es notwendig, dass wir uns neben den Zeilen- auch mit **Spaltenumformungen** einer Matrix befassen. Unter einer **elementaren** Spaltenumformungen verstehen wir, dass die zu den Zeilenumformungen analogen Operationen auf die Spalten einer Matrix $A \in \mathcal{M}_{m \times n, K}$ angewendet werden, also im einzelnen

- (i) die Multiplikation der k -ten Spalten einer Matrix mit einem Wert λ , wobei $\lambda \in K^\times$ und $k \in \{1, \dots, n\}$ ist
- (ii) die Addition des λ -fachen der k -ten Spalte zur ℓ -ten, mit $\lambda \in K$ und $k, \ell \in \{1, \dots, n\}$, $k \neq \ell$.

(7.14) Lemma Die Multiplikationen einer Matrix $A \in \mathcal{M}_{m \times n, K}$ mit den Transponierten von Elementarmatrizen von rechts bewirken elementare Spaltenumformungen. Genauer gilt:

- (i) Die Matrix $A {}^tM_{k, \lambda}$ entsteht aus der Matrix A durch Multiplikation der k -ten Spalte mit λ .
- (ii) Die Matrix $A {}^tA_{k, \ell, \lambda}$ entsteht aus der Matrix A durch Addition des λ -fachen der k -ten Spalte zur ℓ -ten Spalte.

Beweis: Wir beschränken uns auf den Beweis der Aussage (i). Nach der Rechenregel (iv) in Proposition (5.15) gilt $A {}^tM_{k, \lambda} = {}^t({}^tA) {}^tM_{k, \lambda} = {}^t(M_{k, \lambda} {}^tA)$. Der Übergang ${}^tA \mapsto M_{k, \lambda} {}^tA$ bewirkt nach Proposition (7.10) die Multiplikation der k -ten Zeile von tA mit dem Wert λ . Für jedes ℓ ist die ℓ -te Spalte von A gleich der ℓ -ten Zeile von tA , und entsprechend ist die ℓ -te Spalte von $A {}^tM_{k, \lambda} = {}^t(M_{k, \lambda} {}^tA)$ gleich der ℓ -ten Zeile von $M_{k, \lambda} {}^tA$. Also stimmt die ℓ -te Spalte von A mit der ℓ -ten Spalte von $A {}^tM_{k, \lambda}$ für $\ell \neq k$ überein. Für $\ell = k$ unterscheiden sie sich um den Faktor λ . $\square$

Wir bemerken noch, dass mit jeder Matrix $A \in \text{GL}_n(K)$ auch die Transponierte tA invertierbar ist, mit $({}^tA)^{-1} = {}^t(A^{-1})$. Dies folgt direkt aus der Rechnung

$${}^tA {}^t(A^{-1}) = {}^t(A^{-1}A) = {}^tE^{(n)} = E^{(n)}$$

und einer analogen Rechnung, die ${}^t(A^{-1}) {}^tA = E^{(n)}$ liefert.

(7.15) Satz Für jede Matrix $A \in \mathcal{M}_{m \times n, K}$ gibt es invertierbare Matrizen $T \in \text{GL}_m(K)$ und $U \in \text{GL}_n(K)$ und ein $r \in \{1, \dots, n\}$, so dass die Matrix TAU die Blockgestalt

$$\begin{pmatrix} E^{(r)} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix} \text{ besitzt.}$$

Beweis: Wir wissen bereits, dass eine Matrix $T \in \text{GL}_m(K)$ existiert, so dass $B = TA$ in normierter ZSF vorliegt, mit gewissen Kennzahlen $r, j_1, \dots, j_r$. Nach Lemma (7.14) genügt es nun zu zeigen, dass B durch elementare Spaltenumformungen auf die angegebene Blockgestalt gebracht werden kann. Nach Definition der normierten ZSF befindet sich für $1 \leq k \leq r$ in der j_k -ten Spalten von B jeweils der k -te Einheitsvektor $e_k \in K^m$. Nun führt man nacheinander für $1 \leq k \leq r$ die folgende Operation aus:

Addition des $(-a_{k\ell})$ -fachen der j_k -ten Spalte zur ℓ -ten, für $j_k < \ell \leq n$

Durch diese Operation werden die Einträge rechts von der Position (k, j_k) zu Null, während alle übrigen Einträge der Matrix unverändert bleiben.

Nach Durchführung dieser Schritte enthält die modifizierte Matrix B' in den Spalten $j_1, \dots, j_r$ die Einheitsvektoren $e_1, \dots, e_r$, alle übrigen Spalten sind Null. Nun vertauscht man die Spalten noch so, dass sich die Einheitsvektoren in den ersten r Spalten befinden. Dann hat die Matrix die gewünschte Form. $\square$

Auch hier lassen sich die Matrizen T und U , die die angegebene Blockgestalt erzeugen, explizit berechnen. Zunächst wendet man die erforderlichen Zeilenumformungen statt auf A auf die Blockmatrix $(A \ E^{(m)})$ an und erhält so eine Matrix der Form $(B \ T)$ mit $T \in \text{GL}_m(K)$, wobei $B = TA$ sich in normierter ZSF befindet. Anschließend wendet man auf die linke Teilmatrix von $(B \ E^{(n)})$ Spaltenumformungen an, die B auf die Blockgestalt bringen, und **dieselben** Spaltenumformungen auch auf die rechte Teilmatrix. Man erhält damit eine Matrix der Form $(C \ U)$ mit $U \in \text{GL}_m(K)$, wobei $C = BU$ die angegebene Blockgestalt hat. Die Matrizen T und U haben die gewünschte Umformungeigenschaft.

Durch dieses Rechenverfahren haben wir nun auch ein negatives Kriterium für Invertierbarkeit.

(7.16) Satz Sei $A \in \mathcal{M}_{n,K}$ eine Matrix, die durch elementare Zeilenumformungen auf normierte ZSF mit Zeilenrang $r < n$ gebracht werden kann. Dann ist A *nicht* invertierbar.

Beweis: Nehmen wir an, dass die Voraussetzung erfüllt ist, die Matrix A aber dennoch in $\text{GL}_n(K)$ liegt. Nach Satz (7.15) gibt es Matrizen $T, U \in \text{GL}_n(K)$ mit

$$TAU = \begin{pmatrix} E^{(r)} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{pmatrix}.$$

Mit A wäre dann auch TAU invertierbar. Aber eine Matrix mit Nullzeilen kann nicht invertierbar sein, denn für beliebiges $B \in \mathcal{M}_{r \times n, K}$ und $V \in \text{GL}_n(K)$ gilt

$$\begin{pmatrix} B \\ \mathbf{0} \end{pmatrix} V = \begin{pmatrix} BV \\ \mathbf{0}V \end{pmatrix} = \begin{pmatrix} BV \\ \mathbf{0} \end{pmatrix},$$

wobei die letzte Matrix offensichtlich *nicht* mit der Einheitsmatrix übereinstimmt. Der Widerspruch zeigt, dass die Annahme falsch war. $\square$

Unser Rechenverfahren zur Bestimmung der Inversen einer Matrix A liefert also zugleich ein Entscheidungskriterium für die Invertierbarkeit: Kommt bei der Rechnung eine normierte ZSF mit Zeilenrang $r < n$ heraus, dann existiert die Inverse von A nicht.

§ 8. Linearkombinationen und lineare Unabhängigkeit

Inhaltsübersicht

Sei K ein Körper und V ein K -Vektorraum. Ist (v, w) ein Paar bestehend aus Vektoren $v, w \in V$, dann wird jeder Vektor der Form $\lambda v + \mu w$ mit $\lambda, \mu \in K$ eine *Linearkombination* von v und w genannt. Folgt für alle $\lambda, \mu \in K$ aus $\lambda v + \mu w = 0_V$ jeweils $\lambda = \mu = 0_K$, dann bezeichnet man das Paar (v, w) als *linear unabhängig*. Beide Begriffe lassen sich auf beliebig lange Tupel von Vektoren aus V verallgemeinern.

Sei nun $S \subseteq V$ eine beliebige, möglicherweise auch unendlich große, Teilmenge von V . Dann bilden die Linearkombinationen, die mit Vektoren aus der Menge S gebildet werden können, ihrerseits einen Untervektorraum von V , den wir mit $\langle V \rangle_K$ bezeichnen und den *von S erzeugten* Untervektorraum nennen. Es handelt sich um den kleinsten Untervektorraum U von V , der S als Teilmenge enthält. Man bezeichnet die Menge S als *linear unabhängig*, wenn jede Linearkombination eines Tupels von lauter verschiedenen Vektoren aus S linear unabhängig ist. Beide Begriffe werden im folgenden Kapitel bei der Definition des Dimensionsbegriffs eine wichtige Rolle spielen.

Wichtige Begriffe und Sätze

- Linearkombination eines Tupels von Vektoren
- von einer Menge S erzeugter Untervektorraum $\langle S \rangle_K$, Erzeugendensystem eines Untervektorraums
- lineare Unabhängigkeit eines Tupels von Vektoren
- lineare Unabhängigkeit einer Teilmenge S
- Polynomring $R[x]$ über einem Ring R

Im Verlauf dieses Kapitels werden wir es häufig mit Tupeln $(v_1, \dots, v_r)$ bestehend Vektoren, also Elementen eines K -Vektorraums, zu tun haben. Dabei ist es praktisch, auch den Fall $r = 0$ zuzulassen. Das leere Tupel bestehend aus null Vektoren bezeichnen wir mit $()$.

(8.1) Definition Sei V ein K -Vektorraum, $r \in \mathbb{N}_0$ und $(v_1, \dots, v_r)$ ein Tupel von Elementen aus V . Wir bezeichnen einen Vektor $w \in V$ als **Linearkombination** des Tupels, wenn ein Tupel $(\lambda_1, \dots, \lambda_r) \in K^r$ existiert, so dass

$$w = \sum_{i=1}^r \lambda_i v_i \quad \text{erfüllt ist.}$$

Ist (v, w) ein Paar von Vektoren, dann sind die Linearkombinationen von (v, w) genau die Vektoren der Form $\lambda v + \mu w$ mit $\lambda, \mu \in K$. Beispielsweise kann jeder Vektor $v = (a, b) \in \mathbb{R}^2$ als Linearkombination der Einheitsvektoren $e_1 = (1, 0)$ und $e_2 = (0, 1)$ dargestellt werden, denn es ist

$$v = \begin{pmatrix} a \\ b \end{pmatrix} = a \begin{pmatrix} 1 \\ 0 \end{pmatrix} + b \begin{pmatrix} 0 \\ 1 \end{pmatrix} = ae_1 + be_2.$$

Besteht unser Tupel (v) nur aus einem einzigen Vektor, dann sind die Linearkombinationen dieses Tupels genau die Vektoren der Form λv mit $\lambda \in K$, also genau die skalaren Vielfachen von v . Die einzige Linearkombination des leeren Tupels $()$ ist der Nullvektor 0_V , denn jede Summe bestehend aus null Summanden ist definitionsgemäß gleich null.

(8.2) Definition Sei V ein K -Vektorraum und $S \subseteq V$ eine beliebige Teilmenge. Dann bezeichnen wir mit

$$\langle S \rangle_K = \left\{ \sum_{k=1}^r \lambda_k v_k \mid r \in \mathbb{N}_0, \lambda_1, \dots, \lambda_r \in K, v_1, \dots, v_r \in S \right\}$$

die Menge aller Linearkombinationen von Tupeln bestehend aus Vektoren der Menge S . Man nennt $\langle S \rangle_K$ den von S **erzeugten** oder **aufgespannten** Untervektorraum.

Die Rechtfertigung für diese Bezeichnung wird durch folgenden Satz geliefert.

(8.3) Satz Sei V ein K -Vektorraum und $S \subseteq V$ eine Teilmenge. Dann gilt

- (i) Die Menge $\langle S \rangle_K$ bildet einen Untervektorraum von V mit $\langle S \rangle_K \supseteq S$.
- (ii) Ist U ein weiterer Untervektorraum von V mit $U \supseteq S$, dann gilt $U \supseteq \langle S \rangle_K$.

Es handelt sich also bei $\langle S \rangle_K$ um den *kleinsten* Untervektorraum von V , der S als Teilmenge enthält.

Beweis: zu (i) Zunächst beweisen wir, dass $\langle S \rangle_K$ ein Untervektorraum von V ist. Der Nullvektor 0_V ist eine Linearkombination des leeren Tupels $()$ und somit in $\langle S \rangle_K$ enthalten. Seien nun $v, w \in \langle S \rangle_K$ und $\lambda \in K$ vorgegeben. Zu zeigen ist $v + w \in \langle S \rangle_K$ und $\lambda v \in \langle S \rangle_K$. Wegen $v \in \langle S \rangle_K$ gibt es ein $r \in \mathbb{N}_0$ und ein Tupel $(v_1, \dots, v_r)$, so dass v Linearkombination dieses Tupels ist. Es existiert also ein Tupel $(\lambda_1, \dots, \lambda_r) \in K^r$ mit $v = \sum_{i=1}^r \lambda_i v_i$. Genauso folgt aus $w \in \langle S \rangle_K$ die Existenz eines $s \in \mathbb{N}_0$ und von $w_1, \dots, w_s \in S$ und $\mu_1, \dots, \mu_s \in K$ mit $w = \sum_{j=1}^s \mu_j w_j$. Die Gleichung

$$v + w = \sum_{i=1}^r \lambda_i v_i + \sum_{j=1}^s \mu_j w_j$$

zeigt, dass $v + w$ eine Linearkombination des Tupels $(v_1, \dots, v_r, w_1, \dots, w_s)$ ist und somit in $\langle S \rangle_K$ liegt. Ebenso folgt aus $\lambda v = \sum_{i=1}^r (\lambda \lambda_i) v_i$, dass λv eine Linearkombination von $(v_1, \dots, v_r)$ ist und λv somit ebenfalls in $\langle S \rangle_K$ enthalten ist.

Der Nachweis der Untervektorraum-Eigenschaft von $\langle S \rangle_K$ ist damit abgeschlossen. Es gilt $S \subseteq \langle S \rangle_K$, denn jedes $v \in S$ ist wegen $v = 1_K \cdot v$ jeweils Linearkombination des einelementigen Tupels (v) und somit nach Definition in $\langle S \rangle_K$ enthalten.

zu (ii) Sei U ein beliebiger Untervektorraum von V mit $U \supseteq S$. Wir zeigen durch vollständige Induktion über $r \in \mathbb{N}_0$, dass jede Linearkombination jedes r -elementigen Tupels $(v_1, \dots, v_r)$ mit $v_k \in S$ für $1 \leq k \leq r$ in U enthalten ist. Daraus folgt dann unmittelbar $\langle S \rangle_K \subseteq U$. Für $r = 0$ ist die Aussage klar, denn die einzige Linearkombination des leeren Tupels $()$ ist der Nullvektor 0_V , und es gilt $0_V \in U$, weil U ein Untervektorraum von V ist.

Sei nun $r \in \mathbb{N}_0$, und setzen wir die Aussage für dieses r voraus. Sei $(v_1, \dots, v_{r+1})$ ein Tupel mit $v_k \in S$ für $1 \leq k \leq r+1$, und sei w eine Linearkombination dieses Tupels. Es gibt dann also $\lambda_1, \dots, \lambda_{r+1} \in K$ mit $w = \sum_{k=1}^{r+1} \lambda_k v_k$. Der Vektor $w' = \sum_{k=1}^r \lambda_k v_k$ ist eine Linearkombination des r -elementigen Tupels $(v_1, \dots, v_r)$ und somit nach Induktionsvoraussetzung in U enthalten. Weiter gilt $v_{r+1} \in U$ wegen $S \subseteq U$ und weiter $\lambda_{r+1} v_{r+1} \in U$ und $w = w' + \lambda_{r+1} v_{r+1} \in U$, da U ein Untervektorraum von V ist. Damit ist der Induktionsbeweis abgeschlossen. $\square$

Im umgekehrt ein Untervektorraum U eines K -Vektorraums V vorgegeben, dann kann man danach fragen, für welche Teilmengen $S \subseteq V$ jeweils $U = \langle S \rangle_K$ erfüllt ist.

(8.4) Definition Ist V ein K -Vektorraum und $U \subseteq V$ ein Untervektorraum, dann wird jede Teilmenge $S \subseteq V$ mit der Eigenschaft $U = \langle S \rangle_K$ ein **Erzeugendensystem** von U genannt.

Wegen Satz (8.3) erfüllt jedes Erzeugendensystem S von U notwendigerweise die Bedingung $S \subseteq U$, und insbesondere gilt $\langle U \rangle_K = U$. In der Regel ist die Teilmenge S im Vergleich zu U eine sehr kleine Menge und besteht häufig nur aus endlich vielen Elementen.

Im nächsten Kapitel werden wir uns oft auch für Erzeugendensysteme des gesamten Vektorraums V , also Teilmengen $S \subseteq V$ mit $\langle S \rangle_K = V$, interessieren. Für die Vektorräume K^n und $\mathcal{M}_{m \times n, K}$ existieren endliche Erzeugendensysteme, nämlich im ersten Fall die Menge $\{e_1, \dots, e_n\}$ der Einheitsvektoren und im zweiten die Menge $\{B_{k\ell} \mid 1 \leq k \leq m, 1 \leq \ell \leq n\}$ der Basismatrizen. Es gibt aber auch Vektorräume, die durch keine endliche Teilmenge ihrer Vektoren aufgespannt werden können. Um zumindest ein konkretes Beispiel für eine solche Situation vor Augen zu haben, und weil wir ihn auch in den späteren Kapiteln noch brauchen werden, führen wir den Polynomring über einem Ring R ein.

(8.5) Satz (ohne Beweis)

Für jeden Ring R gibt es einen Erweiterungsring $R[x] \supseteq R$ mit einem ausgezeichneten Element $x \notin R$, so dass für jedes $f \in R[x]$ folgende Bedingung erfüllt ist. Entweder ist $f = 0_R$, oder es gibt ein eindeutig bestimmtes $n \in \mathbb{N}_0$ und eindeutig bestimmte Elemente $a_0, \dots, a_n \in R$ mit

$$f = \sum_{k=0}^n a_k x^k \quad \text{und} \quad a_n \neq 0_R.$$

Man nennt $R[x]$ den **Polynomring** über dem Körper R , seine Elemente bezeichnet man als **Polynome**. Im Fall $f \neq 0_R$ bezeichnet man n als den **Grad** des Polynoms.

Jedem Element $f = \sum_{k=0}^n a_k x^k \in R[x]$ kann eine Abbildung $R \rightarrow R$ zugeordnet werden, die ein Element $c \in R$ jeweils auf

$$f(c) = \sum_{k=0}^n a_k c^k$$

abbildet. Man nennt diese Abbildung die dem Polynom f zugeordnete **Polynomfunktion**. Ein häufiger Fehler besteht darin, ein Polynom f mit seiner Polynomfunktion gleichzusetzen. Diese Gleichsetzung ist aber unzulässig und kann zu Widersprüchen führen. Der Grund dafür ist, dass bei einem endlichen Ring R verschiedene Polynome dieselbe Polynomfunktion liefern können. Ist beispielsweise $R = \mathbb{F}_2 = \{\bar{0}, \bar{1}\}$, der Körper mit zwei Elementen, dann haben die Polynome $x \in \mathbb{F}_2[x]$ und $x^2 \in \mathbb{F}_2[x]$ dieselbe zugeordnete Polynomfunktion, gegeben durch $\bar{0} \mapsto \bar{0}$ und $\bar{1} \mapsto \bar{1}$.

In diesem Zusammenhang sei auch darauf hingewiesen, dass das Element x des Polynomrings $R[x]$ kein Element des Rings R ist! Statt einer Funktion $R \rightarrow R$ sollte man sich unter einem Polynom in $R[x]$ deshalb eher einen rein „formalen Ausdruck“, gebildet mit Elementen des Rings R und der Variablen x des Polynomrings, vorstellen.

Nach diesem kurzen Exkurs kehren wir nun zur Linearen Algebra zurück. Ist K ein Körper, dann kann man $K[x]$ als K -Vektorraum betrachten, in dem man die Vektoraddition mit der gewöhnlichen Addition im Ring $K[x]$ gleichsetzt und die skalare Multiplikation $K \times K[x] \rightarrow K[x]$ durch Einschränkung der Multiplikationsabbildung

$$K[x] \times K[x] \longrightarrow K[x] \quad , \quad (f, g) \mapsto fg$$

definiert. In diesem Fall ist dann $S = \{x^n \mid n \in \mathbb{N}_0\}$ ein (unendliches) Erzeugendensystem von $K[x]$ als K -Vektorraum, wobei $x^0 = 1_K$ gesetzt wird. Jedes Polynom ist Linearkombination von S . Beispielsweise ist das Polynom $f = x^7 - 4x^3 + 5$ eine Linearkombination des Tupels $(1, x^3, x^7)$, und ebenso eine Linearkombination von $(1, x, x^2, \dots, x^7)$.

Man kann sich aber leicht überzeugen, dass keine *endliche* Teilmenge $T \subseteq K[x]$ mit der Eigenschaft $\langle T \rangle_K = K[x]$ existiert. Denn in T gibt es ein Polynom mit maximalem Grad n , und folglich kann keine Linearkombination von T einen Grad größer als n haben. Dies bedeutet, dass zum Beispiel x^{n+1} nicht in $\langle T \rangle_K$ enthalten ist.

Kommen wir nun zur zweiten wichtigen Definition dieses Kapitels.

(8.6) Definition Sei V ein K -Vektorraum. Wir bezeichnen ein Tupel $(v_1, \dots, v_r)$ mit $r \in \mathbb{N}_0$, bestehend aus Vektoren $v_k \in V$, als **linear unabhängig**, wenn für jedes Tupel $(\lambda_1, \dots, \lambda_r) \in K^r$ die Implikation

$$\sum_{k=1}^r \lambda_k v_k = 0_V \quad \Rightarrow \quad \lambda_1 = \dots = \lambda_r = 0_K \quad \text{erfüllt ist.}$$

erfüllt ist. Ein Tupel, das nicht linear unabhängig ist, wird **linear abhängig** genannt.

Das leere Tupel $()$ ist nach Definition linear unabhängig. Ein Tupel (v) bestehend aus einem einzelnen Vektor v ist genau dann linear unabhängig, wenn $v \neq 0_V$ ist. Ein Paar (v, w) bestehend aus zwei Vektoren ist genau dann linear unabhängig, wenn $v \neq 0_V$ und w kein skalares Vielfaches von v ist, also $w \neq \lambda v$ für alle $\lambda \in K$ gilt.

Für jedes $n \in \mathbb{N}$ ist im K -Vektorraum K^n das Tupel $(e_1, \dots, e_n)$ bestehend aus den Einheitsvektoren linear unabhängig. Sei nämlich $(\lambda_1, \dots, \lambda_n) \in K^n$ ein Tupel mit $\sum_{k=1}^n \lambda_k v_k = 0_{K^n}$. Für $1 \leq i \leq n$ ist die i -te Komponente von $\sum_{k=1}^n \lambda_k e_k$ jeweils gleich $\sum_{k=1}^n \lambda_k \delta_{ki} = \lambda_i$. Aus $\sum_{k=1}^n \lambda_k v_k = 0_{K^n}$ folgt also $\lambda_i = 0_K$ für $1 \leq i \leq n$. Also ist $(e_1, \dots, e_n)$ linear unabhängig.

Auch für den Begriff der linearen Unabhängigkeit gibt es eine Variante, die sich auf Mengen bezieht.

(8.7) Definition Sei V ein K -Vektorraum. Eine Teilmenge $S \subseteq V$ bezeichnen wir als **linear unabhängig**, wenn jedes Tupel $(v_1, \dots, v_r)$ bestehend aus lauter verschiedenen Elementen v_k der Menge S linear unabhängig ist.

Die Einschränkung in der Definition ist sinnvoll, denn ein Tupel, indem derselbe Vektor mehrfach vorkommt, ist automatisch linear abhängig. Sei nämlich $r \in \mathbb{N}$, $r \geq 2$ und $(v_1, \dots, v_r)$ ein Tupel bestehend aus Elementen von V . Nehmen wir weiter an, dass es natürliche Zahlen i, j mit $1 \leq i < j \leq r$ und $v_i = v_j$ gibt. Definieren wir dann $\lambda_i = 1_K$, $\lambda_j = -1_K$ und $\lambda_k = 0_K$ für alle $k \in \{1, \dots, r\} \setminus \{i, j\}$, dann gilt $\sum_{k=1}^r \lambda_k v_k = 0_V$, ohne dass die Koeffizienten $\lambda_1, \dots, \lambda_r$ alle gleich null sind.

Wir notieren einige elementare Beobachtungen, die sich direkt aus der Definition ergeben.

(8.8) Lemma Sei V ein K -Vektorraum.

- (i) Ist $S \subseteq V$ linear unabhängig, dann gibt dasselbe für jede Teilmenge $T \subseteq S$. Insbesondere ist die leere Menge stets linear unabhängig.
- (ii) Jede Teilmenge $S \subseteq V$ mit $0_V \in S$ ist linear abhängig.
- (iii) Ist $r \in \mathbb{N}_0$ und sind $v_1, \dots, v_r \in V$ mit $v_i \neq v_j$ für $i \neq j$, so ist r -elementige Menge $\{v_1, \dots, v_r\}$ genau dann linear unabhängig, wenn das Tupel $(v_1, \dots, v_r)$ linear unabhängig ist.

Betrachten wir noch einige konkrete Beispiele.

- (i) Die Menge $\{e_1, \dots, e_n\}$ der Einheitsvektoren ist eine n -elementige linear unabhängige Teilmenge des K -Vektorraums K^n . Dies ergibt sich aus Teil (iii) des soeben formulierten Lemmas in Verbindung mit der Feststellung von oben, dass das n -Tupel $(e_1, \dots, e_n)$ linear unabhängig ist.
- (ii) Auf ähnliche Weise kann man zeigen, dass die Menge $\{B_{k\ell}^{(m \times n)} \mid 1 \leq k \leq m, 1 \leq \ell \leq n\}$ der Basismatrizen im K -Vektorraum $\mathcal{M}_{m \times n, K}$ der $m \times n$ -Matrizen linear unabhängig ist.
- (iii) In einigen Vektorräumen existieren auch unendliche linear unabhängige Mengen. Beispielsweise ist im Polynomring $K[x]$, aufgefasst als K -Vektorraum, die Teilmenge $\{x^n \mid n \in \mathbb{N}_0\}$ bestehend aus den sog. *Monomen*, linear unabhängig.

(8.9) Proposition Sei V ein K -Vektorraum und $S \subseteq V$ eine beliebige Teilmenge.

- (i) Genau dann ist S linear abhängig, wenn ein $v \in S$ mit $v \in \langle S \setminus \{v\} \rangle_K$ existiert.
- (ii) Sie ist genau dann linear unabhängig, wenn $v \notin \langle S \setminus \{v\} \rangle_K$ für alle $v \in S$ erfüllt ist.
- (iii) Ist S linear unabhängig und $v \in V \setminus \langle S \rangle_K$, dann ist auch $S \cup \{v\}$ linear unabhängig.

Beweis: zu (i) „ $\Leftarrow$ “ Angenommen, es gibt ein $v \in S$ mit $v \in \langle S \setminus \{v\} \rangle_K$. Dann gibt es ein $r \in \mathbb{N}_0$ und ein Tupel $(v_1, \dots, v_r)$ von Vektoren aus $S \setminus \{v\}$ mit der Eigenschaft, dass v eine Linearkombination dieses Tupels ist. Es gibt also ein Tupel $(\lambda_1, \dots, \lambda_r) \in K^r$ mit $v = \sum_{k=1}^r \lambda_k v_k$. Dabei dürfen wir davon ausgehen, dass die Vektoren $v_1, \dots, v_r$ alle verschieden sind. Kommt nämlich einer der Vektoren v_i mehrfach im Tupel vor, gilt also $v_j = v_i$ für ein $j \neq i$, dann können wir die Summe $\lambda_i v_i + \lambda_j v_j$ durch $(\lambda_i + \lambda_j) v_i$ ersetzen. Dies zeigt, dass wir v_j aus dem Tupel streichen können, ohne dass die Eigenschaft von v , eine Linearkombination des Tupels zu sein, verloren geht. Die Gleichung

$$v = \sum_{k=1}^r \lambda_k v_k \iff (-1_K)v + \sum_{i=1}^r \lambda_i v_i = 0_V$$

zeigt nun, dass das Tupel $(v, v_1, \dots, v_r)$ linear abhängig ist. Weil das Tupel aus lauter verschiedenen Elementen der Menge S besteht, folgt daraus, dass S linear abhängig ist.

„ $\Rightarrow$ “ Ist S linear abhängig, dann gibt es ein $r \in \mathbb{N}$ und ein linear abhängiges Tupel $(v_1, \dots, v_r)$ bestehend aus lauter verschiedenen Vektoren der Menge S . Die lineare Abhängigkeit bedeutet, dass $\lambda_1, \dots, \lambda_r \in K$ existieren, nicht alle gleich Null, mit $\sum_{k=1}^r \lambda_k v_k = 0_V$. Nehmen wir an, dass $i \in \{1, \dots, r\}$ ein Index mit $\lambda_i \neq 0$ ist. Dann kann die Gleichung umgestellt werden zu

$$v_i = \sum_{j \neq i} \left(-\frac{\lambda_j}{\lambda_i} \right) v_j.$$

Dies zeigt, dass v_i in $\langle S \setminus \{v_i\} \rangle_K$ enthalten ist.

Die Aussage (ii) folgt rein logisch aus Teil (i) durch Negation auf beiden Seiten der Äquivalenz.

zu (iii) Nehmen wir an, dass S linear unabhängig ist, dass $v \notin \langle S \rangle_K$ gilt, und dass $S \cup \{v\}$ dennoch linear abhängig ist. Dann gibt es ein $r \in \mathbb{N}$ und ein linear abhängiges Tupel $(v_1, \dots, v_r)$ bestehend aus lauter verschiedenen Elementen der Menge $S \cup \{v\}$. Einer dieser Vektoren v_i muss mit v übereinstimmen, denn ansonsten wäre $(v_1, \dots, v_r)$ ein linear abhängiges Tupel bestehend aus Elementen der Menge S , und S somit linear abhängig, im Widerspruch zur Voraussetzung.

Sei $i \in \{1, \dots, r\}$ der Index mit $v_i = v$. Auf Grund der linearen Abhängigkeit gibt es Koeffizienten $\lambda_1, \dots, \lambda_r \in K$, nicht alle gleich Null, mit $\sum_{k=1}^r \lambda_k v_k = 0_V$. Dabei muss $\lambda_i \neq 0_K$ gelten, denn andernfalls wäre das Tupel ohne den Vektor v_i ebenfalls linear abhängig, was erneut im Widerspruch zur linearen Unabhängigkeit von S stehen würde. So aber können wir die Gleichung wieder zu

$$v = v_i = \sum_{j \neq i} \left(-\frac{\lambda_j}{\lambda_i} \right) v_j$$

umstellen. Aber dies steht im Widerspruch zur Voraussetzung $v \notin \langle S \rangle_K$. □

Am Ende dieses Kapitels sehen wir uns an, wie sich lineare Unabhängigkeit und die Existenz von Linearkombinationen rechnerisch nachweisen lässt. Sei $V = \mathbb{R}^3$ und (v_1, v_2, v_3) das Tupel bestehend aus den drei Vektoren

$$v_1 = \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 2 \\ 6 \\ 7 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 0 \\ 8 \\ 7 \end{pmatrix}.$$

Sei außerdem $v = (1, 0, 1)$ und $w = (3, 5, 7)$. Unser Ziel ist es zu überprüfen, ob v bzw. w Linearkombinationen des Tupels (v_1, v_2, v_3) sind. Dass es sich bei v um eine Linearkombination des Tupels handelt, ist äquivalent zur Existenz von Koeffizienten $\lambda_1, \lambda_2, \lambda_3 \in \mathbb{R}$ mit

$$\begin{aligned} \lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = v &\Leftrightarrow \lambda_1 \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 2 \\ 6 \\ 7 \end{pmatrix} + \lambda_3 \begin{pmatrix} 0 \\ 8 \\ 7 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \Leftrightarrow \\ &\begin{pmatrix} \lambda_1 + 2\lambda_2 \\ -\lambda_1 + 6\lambda_2 + 8\lambda_3 \\ 7\lambda_2 + 7\lambda_3 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \end{aligned}$$

$$\Leftrightarrow (\lambda_1, \lambda_2, \lambda_3) \text{ ist Lösungsmenge des LGS } x_1 + 2x_2 = 1, -x_1 + 6x_2 + 8x_3 = 0, 7x_2 + 7x_3 = 1.$$

Um zu sehen, ob das LGS eine Lösung hat, stellen wir die erweiterte Koeffizientenmatrix auf und bringen sie auf normierte Zeilenstufenform.

$$\begin{pmatrix} 1 & 2 & 0 & 1 \\ -1 & 6 & 8 & 0 \\ 0 & 7 & 7 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 1 \\ 0 & 8 & 8 & 1 \\ 0 & 7 & 7 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 1 \\ 0 & 1 & 1 & 0 \\ 0 & 7 & 7 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 1 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & -2 & 1 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Die dritte Zeile in der letzten Matrix entspricht der Gleichung $0 = 1$. Das LGS ist also *unlösbar*. Auf Grund der oben formulierten Äquivalenz folgt daraus, dass v keine Linearkombination des Tupels (v_1, v_2, v_3) ist.

Betrachten wir nun an Stelle von v den Vektor w . Hier führt die Gleichung $\lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = w$ nach dem gleichen Schema auf das LGS

$$x_1 + 2x_2 = 3, -x_1 + 6x_2 + 8x_3 = 5, 7x_2 + 7x_3 = 7.$$

Wieder stellen wir die erweiterte Koeffizientenmatrix auf und formen auf normierte ZSF um.

$$\begin{pmatrix} 1 & 2 & 0 & 3 \\ -1 & 6 & 8 & 5 \\ 0 & 7 & 7 & 7 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 3 \\ 0 & 8 & 8 & 8 \\ 0 & 7 & 7 & 7 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 3 \\ 0 & 1 & 1 & 1 \\ 0 & 1 & 1 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 2 & 0 & 3 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & -2 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Die Gleichungen in der letzten Matrix lauten $x_1 - 2x_3 = 1$ und $x_2 + x_3 = 1$, was zu $x_1 = 1 + 2x_3$ und $x_2 = 1 - x_3$ umgeformt werden kann. Die Lösungsmenge $\mathcal{L}$ des ursprünglichen LGS ist also gegeben durch

$$\mathcal{L} = \left\{ \begin{pmatrix} 1 + 2x_3 \\ 1 - x_3 \\ x_3 \end{pmatrix} \mid x_3 \in \mathbb{R} \right\} = \left\{ \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix} + \lambda \begin{pmatrix} 2 \\ -1 \\ 1 \end{pmatrix} \mid \lambda \in \mathbb{R} \right\}.$$

Jedes Element der Lösungsmenge liefert eine Darstellung von w als Linearkombination des Tupels (v_1, v_2, v_3) . Setzt man $\lambda = 0$, so erhält man zum Beispiel das Element $(1, 1, 0) \in \mathcal{L}$, und $\lambda = 1$ entspricht $(3, 0, 1) \in \mathcal{L}$. Tatsächlich gilt sowohl

$$1 \cdot \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} + 1 \cdot \begin{pmatrix} 2 \\ 6 \\ 7 \end{pmatrix} + 0 \cdot \begin{pmatrix} 0 \\ 8 \\ 7 \end{pmatrix} = \begin{pmatrix} 3 \\ 5 \\ 7 \end{pmatrix} \quad \text{als auch} \quad 3 \cdot \begin{pmatrix} 1 \\ -1 \\ 0 \end{pmatrix} + 0 \cdot \begin{pmatrix} 2 \\ 6 \\ 7 \end{pmatrix} + 1 \cdot \begin{pmatrix} 0 \\ 8 \\ 7 \end{pmatrix} = \begin{pmatrix} 3 \\ 5 \\ 7 \end{pmatrix}.$$

Nach einem ähnlichen Schema lässt sich auch die lineare Unabhängigkeit behandeln. Diesmal betrachten wir in $V = \mathbb{R}^3$ das Tupel (v_1, v_2, v_3) bestehend aus den Vektoren

$$v_1 = \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 0 \\ 2 \\ -1 \end{pmatrix}, \quad v_3 = \begin{pmatrix} 3 \\ 3 \\ 3 \end{pmatrix}.$$

Diesmal besteht unser Ziel darin, die lineare Unabhängigkeit von (v_1, v_2, v_3) nachzuweisen. Für jedes Tripel $(\lambda_1, \lambda_2, \lambda_3)$ in $\mathbb{R}^3$ gilt die Äquivalenz

$$\lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = 0_{\mathbb{R}^3} \iff \lambda_1 \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} + \lambda_2 \begin{pmatrix} 0 \\ 2 \\ -1 \end{pmatrix} + \lambda_3 \begin{pmatrix} 3 \\ 3 \\ 3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix} \iff \begin{pmatrix} \lambda_1 + 3\lambda_3 \\ 2\lambda_1 + 2\lambda_2 + 3\lambda_3 \\ -\lambda_2 + 3\lambda_3 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}$$

$$\iff (\lambda_1, \lambda_2, \lambda_3) \text{ ist Lösung des LGS } x_1 + 3x_3 = 0, \quad 2x_1 + 2x_2 + 3x_3 = 0, \quad -x_2 + 3x_3 = 0.$$

Diesmal handelt es sich um ein *homogenes* LGS. Es genügt also, die nicht-erweiterte Koeffizientenmatrix aufzustellen und auf normierte ZSF zu bringen.

$$\begin{pmatrix} 1 & 0 & 3 \\ 2 & 2 & 3 \\ 0 & -1 & 3 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 3 \\ 0 & 2 & -3 \\ 0 & 1 & -3 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 3 \\ 0 & 1 & -3 \\ 0 & 2 & -3 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 3 \\ 0 & 1 & -3 \\ 0 & 0 & 3 \end{pmatrix} \\ \mapsto \begin{pmatrix} 1 & 0 & 3 \\ 0 & 1 & -3 \\ 0 & 0 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Die letzte Matrix entspricht dem System $x_1 = 0$, $x_2 = 0$, $x_3 = 0$, die Lösungsmenge $\mathcal{L}$ des Systems ist also gleich $\{(0, 0, 0)\}$. Daraus ergibt sich insgesamt die Äquivalenz

$$\lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = 0_{\mathbb{R}^3} \iff (\lambda_1, \lambda_2, \lambda_3) = (0, 0, 0) \iff \lambda_1 = \lambda_2 = \lambda_3 = 0.$$

Dies zeigt, dass das Tupel (v_1, v_2, v_3) tatsächlich linear unabhängig ist; dafür ist bereits die Gültigkeit der Implikation „ $\Rightarrow$ “ hinreichend. Das Beispiel zeigt, auf welche Weise die lineare Unabhängigkeit eines Tupels von Vektoren im K^m allgemein überprüft werden kann.

(8.10) Proposition Seien $m, n \in \mathbb{N}$, und sei $(v_1, \dots, v_n)$ ein Tupel bestehend aus Vektoren $v_k \in K^m$. Sei $A \in \mathcal{M}_{m \times n, K}$ die Matrix mit der Eigenschaft, dass v_k jeweils der k -te Spaltenvektor $a_{\bullet k}$ von A ist, für $1 \leq k \leq n$. Genau dann ist Tupel linear unabhängig, wenn die einzige Lösung des homogenen linearen Gleichungssystems $Ax = 0_{K^m}$ der Nullvektor 0_{K^n} ist.

Beweis: Für beliebige $\lambda_1, \dots, \lambda_n \in K$ ist der Vektor $\sum_{k=1}^n \lambda_k v_k \in K^m$ gleich dem Matrix-Vektor-Produkt Aw , mit $w = (\lambda_1, \dots, \lambda_n)$, denn für $1 \leq i \leq m$ ist der i -te Eintrag von Aw nach Definition des Matrix-Vektor-Produkts gleich $\sum_{k=1}^n a_{ik} \lambda_k$, und wegen $v_k = a_{\bullet k} = (a_{1k}, \dots, a_{mk})$ für $1 \leq k \leq n$ ist dies zugleich der i -te Eintrag des Vektors $\sum_{k=1}^n \lambda_k v_k$.

Das Tupel $(v_1, \dots, v_n)$ ist nach Definition genau dann linear unabhängig, wenn das einzige Tupel $(\lambda_1, \dots, \lambda_n)$ mit der Eigenschaft $\sum_{k=1}^n \lambda_k v_k = 0_{K^m}$ der Nullvektor 0_{K^n} ist. Dies wiederum ist auf Grund unserer Vorüberlegung genau dann der Fall, wenn die Gleichung $Aw = 0_{K^m}$ nur von $w = 0_{K^n}$ erfüllt ist, also genau dann, wenn 0_{K^n} die einzige Lösung des linearen Gleichungssystems $Ax = 0_{K^m}$ ist. $\square$

§ 9. Basen eines Vektorraums, Dimensionsbegriff

Inhaltsübersicht

Eine Teilmenge eines K -Vektorraums V wird *Basis* genannt, wenn sie zugleich linear unabhängig und ein Erzeugendensystem von V ist. Das wichtigste Ergebnis dieses Abschnitts lautet, dass je zwei Basen eines endlich erzeugten K -Vektorraums dieselbe Elementezahl besitzen; diese wird dann die *Dimension* von V genannt.

Wir bestimmen Basen für eine Reihe konkreter K -Vektorräume (der K^n , Matrizen, Polynome). Außerdem leiten wir einige fundamentale Aussagen über K -Vektorräume her. Als wichtigste sind hier der *Basisauswahlsatz* und der *Basisergänzungssatz* zu nennen: Aus jedem Erzeugendensystem eines K -Vektorraums kann eine Basis ausgewählt werden, und jede linear unabhängige Teilmenge kann zu einer Basis ergänzt werden.

Wichtige Begriffe und Sätze

- Austauschsatz
- Je zwei Basen eines endlich erzeugten K -Vektorraums haben gleich viele Elemente.
- Basisauswahlsatz
- Basisergänzungssatz

(9.1) Definition Sei V ein K -Vektorraum. Eine Teilmenge $B \subseteq V$ heißt **Basis** von V , wenn sie linear unabhängig und ein Erzeugendensystem von V ist. Ein Tupel $(v_1, \dots, v_n)$ mit $n \in \mathbb{N}_0$, bestehend aus Vektoren $v_i \in V$, wird **geordnete Basis** genannt, wenn $\{v_1, \dots, v_n\}$ aus n verschiedenen Elementen besteht und eine Basis von V bildet.

Im letzten Kapitel haben wir mehrere Beispiele für Teilmengen eines Vektorraums V gesehen, die einerseits den Vektorraum V erzeugen und andererseits auch linear unabhängig sind. Beispielsweise ist durch die Menge $\{e_1, \dots, e_n\}$ der Einheitsvektoren eine Basis des K^n gegeben. Die Menge $\{B_{k\ell}^{(m \times n)} \mid 1 \leq k \leq m, 1 \leq \ell \leq n\}$ der Basismatrizen bildet eine Basis des Vektorraums $\mathcal{M}_{m \times n, K}$ (daher der Name). Die Menge $\{x^n \mid n \in \mathbb{N}_0\}$ bildet eine (unendliche) Basis des K -Vektorraums $K[x]$.

(9.2) Satz Sei V ein K -Vektorraum. Für eine Teilmenge $B \subseteq V$ sind die folgenden Aussagen äquivalent.

- (i) Sie ist eine Basis von V .
- (ii) Sie ist ein minimales Erzeugendensystem von V .
- (iii) Sie ist eine maximale linear unabhängige Teilmenge von V .

Beweis: „(i) $\Rightarrow$ (ii)“ Nehmen wir an, dass B kein minimales Erzeugendensystem von V ist. Dann gibt es eine Teilmenge $S \subsetneq B$ mit $V = \langle S \rangle_K$. Wählen wir $v \in B \setminus S$ beliebig, dann gilt $v \in \langle S \rangle_K$. Nach Prop. (8.9) (i) ist $S \cup \{v\}$ also linear abhängig. Wegen $B \supseteq S \cup \{v\}$ ist dann auch B linear abhängig, im Widerspruch zur Voraussetzung.

„(ii) $\Rightarrow$ (iii)“ Gehen wir zunächst davon aus, dass B linear abhängig ist. Dann gibt es nach Prop. (8.9) (i) ein $v \in B$ mit $v \in \langle B \setminus \{v\} \rangle_K$. Aus $B \setminus \{v\} \subseteq \langle B \setminus \{v\} \rangle_K$ und $v \in \langle B \setminus \{v\} \rangle_K$ folgt $B \subseteq \langle B \setminus \{v\} \rangle_K$. Mit Satz (8.3) folgt $V = \langle B \rangle_K \subseteq \langle B \setminus \{v\} \rangle_K$, weil $\langle B \setminus \{v\} \rangle_K$ ein Untervektorraum von V ist, und damit $V = \langle B \setminus \{v\} \rangle_K$. Aber dies steht im Widerspruch zu der Voraussetzung, dass B ein *minimales* Erzeugendensystem von V ist.

Nehmen wir nun an, B ist zwar linear unabhängig, aber als linear unabhängige Teilmenge nicht maximal. Dann gibt es eine linear unabhängige Teilmenge S von V mit $S \supsetneq B$. Sei nun v ein beliebiges Element in $S \setminus B$. Wegen $V = \langle B \rangle_K$ gilt $v \in \langle B \rangle_K$ und wegen $B \subseteq S \setminus \{v\}$ damit erst recht $v \in \langle S \setminus \{v\} \rangle_K$. Nach Prop. (8.9) (i) bedeutet dies, dass S linear abhängig ist. Unsere Annahme hat also erneut zu einem Widerspruch geführt.

„(iii) $\Rightarrow$ (i)“ Nehmen wir an, dass B kein Erzeugendensystem von V ist. Dann existiert ein $v \in V \setminus \langle B \rangle_K$. Nach Prop. (8.9) (ii) ist deshalb mit B auch $B \cup \{v\}$ linear unabhängig. Aber dies widerspricht der Voraussetzung, dass B maximal als linear unabhängige Teilmenge von V ist. Also ist B sowohl linear unabhängig als auch ein Erzeugendensystem von V , insgesamt eine Basis. $\square$

(9.3) Proposition Sei V ein K -Vektorraum, E ein Erzeugendensystem von V und $B \subseteq E$ eine maximale linear unabhängige Teilmenge von E . Dann ist B eine Basis von V .

Beweis: Nach Voraussetzung ist $B \cup \{v\}$ für jedes $v \in E \setminus B$ linear abhängig. Nach Prop. (8.9) (i) folgt daraus jeweils $v \in \langle B \rangle_K$, es gilt also $E \setminus B \subseteq \langle B \rangle_K$. Zusammen mit $B \subseteq \langle B \rangle_K$ erhalten wir $E \subseteq \langle B \rangle_K$. Weil $\langle B \rangle_K$ ein Untervektorraum von V ist, folgt $V = \langle E \rangle_K \subseteq \langle B \rangle_K$ nach Satz (8.3). Es gilt somit $V = \langle B \rangle_K$, und B ist linear unabhängig. Also ist B insgesamt eine Basis von V . $\square$

(9.4) Lemma (*Austauschlemma*)

Sei V ein K -Vektorraum, $S \subseteq V$ eine linear unabhängige Teilmenge und $E \subseteq V$ ein Erzeugendensystem von V . Dann gibt es für jeden Vektor $v \in S \setminus E$ ein $w \in E \setminus S$, so dass auch $(S \setminus \{v\}) \cup \{w\}$ eine linear unabhängige Teilmenge von V ist.

Beweis: Sei $v \in S \setminus E$ und $S' = S \setminus \{v\}$. Nehmen wir an, dass kein w mit der angegebenen Eigenschaft existiert. Dann ist $(S \setminus \{v\}) \cup \{w\}$ für alle $w \in E \setminus S$ linear abhängig. Nach Prop. (8.9) (ii) folgt daraus $w \in \langle S' \rangle_K$ für alle $w \in E \setminus S$, also $E \setminus S \subseteq \langle S' \rangle_K$. Da auch $S' \subseteq \langle S' \rangle_K$ gilt, ist insgesamt $S' \cup (E \setminus S) = E \setminus \{v\}$ eine Teilmenge von $\langle S' \rangle_K$. Wegen $v \notin E$ gilt außerdem $E \setminus \{v\} = E$. Aus $E \subseteq \langle S' \rangle_K$ folgt nach Satz (8.3), dass $V = \langle E \rangle_K$ in $\langle S' \rangle_K$ enthalten ist. Es gilt also $V = \langle S' \rangle_K$, und insbesondere $v \in \langle S' \rangle_K$. Aber nach Prop. (8.9) (i) ist dann $S = S' \cup \{v\}$ linear abhängig, im Widerspruch zur Voraussetzung. $\square$

(9.5) Satz (Austauschsatz)

Sei V ein K -Vektorraum, $S \subseteq V$ eine linear unabhängige Teilmenge und $E \subseteq V$ ein Erzeugendensystem. Dann gibt es für jede endliche Teilmenge $T \subseteq S$ eine Teilmenge $F \subseteq E$ mit der Eigenschaft, dass $|F| = |T|$ gilt und auch $(S \setminus T) \cup F$ linear unabhängig ist.

Beweis: Der Hauptteil des Beweises besteht im Nachweis der folgenden Hilfsaussage: Für jede endliche Teilmenge $T \subseteq S \setminus E$ gibt es eine Teilmenge $F \subseteq E \setminus S$, so dass $|F| = |T|$ gilt und $(S \setminus T) \cup F$ linear unabhängig ist. Wir beweisen diese Aussage durch vollständige Induktion über $n = |T|$. Ist $n = 0$, dann folgt $T = \emptyset$. Setzen wir $F = \emptyset$, dann gilt $|F| = 0 = |T|$. Außerdem gilt $(S \setminus T) \cup F = S$, also ist diese Menge linear unabhängig.

Sei nun $n \in \mathbb{N}_0$, und setzen wir die Aussage für n voraus. Sei $T \subseteq S \setminus E$ eine $(n+1)$ -elementige Teilmenge und $v \in T$ ein beliebiges Element. Wir setzen $T' = T \setminus \{v\}$. Wegen $|T'| = n$ und $T' \subseteq S \setminus E$ dürfen wir die Induktionsvoraussetzung anwenden und erhalten eine Teilmenge $F' \subseteq E \setminus S$ mit $|F'| = n$ und der Eigenschaft, dass die Menge

$$S' = (S \setminus T') \cup F'$$

linear unabhängig ist. Nun wenden wir das Austauschlemma, Lemma (9.4), auf die linear unabhängige Menge S' , das Erzeugendensystem E und den Vektor v an. Es gilt $v \in S' \setminus E$ (wird noch überprüft, siehe unten); deshalb ist die Anwendung des Austauschlemmas zulässig, und wir erhalten ein $w \in E \setminus S'$ derart, dass $(S' \setminus \{v\}) \cup \{w\}$ linear unabhängig ist.

Wir überprüfen nun, dass die Menge $F = F' \cup \{w\}$ alle gewünschten Eigenschaften hat. Wegen $w \notin F'$ (s.u.) gilt $|F| = n + 1$. Aus $F' \subseteq E \setminus S$ und $w \in E \setminus S$ (s.u.) folgt $F \subseteq E \setminus S$. Darüber hinaus gilt

$$(S' \setminus \{v\}) \cup \{w\} = (((S \setminus T') \cup F') \setminus \{v\}) \cup \{w\} = ((S \setminus T') \setminus \{v\}) \cup F' \cup \{w\} = (S \setminus T) \cup F,$$

wobei die zweite Gleichung dadurch zu Stande kommt, dass $v \notin F'$ gilt (s.u.) und es deshalb gleichgültig ist, ob wir erst F' hinzunehmen und dann v aus der Menge entfernen, oder umgekehrt. Um den Beweis der Hilfsaussage abzuschließen, müssen wir nur noch überprüfen, dass tatsächlich

$$(i) v \in S' \setminus E \quad (ii) w \in E \setminus S \quad (iii) v \notin F' \quad (iv) w \notin F' \quad \text{erfüllt ist.}$$

zu (i) Wegen $v \in S \setminus T'$ gilt $v \in S'$, und aus $v \in T$ und $T \subseteq S \setminus E$ folgt $v \notin E$.

zu (ii) Aus $w \in E \setminus S'$ folgt $w \in E$. Angenommen, es gilt auch $w \in S$. Weil aus $w \notin S'$ insbesondere $w \notin S \setminus T'$ folgt, ist dies nur möglich, wenn w in T' liegt. Aber dies ist wegen $w \in E$ und $T' \subseteq T \subseteq S \setminus E$ nicht der Fall.

zu (iii) Wegen $F' \subseteq E$, $v \in T$ und $T \subseteq S \setminus E$ ist $v \in F'$ ausgeschlossen.

zu (iv) Aus $w \in F'$ würde $w \in S'$ folgen, im Widerspruch zu $w \in E \setminus S'$.

Um nun die eigentliche Aussage des Austauschsatzes zu beweisen, sei $T \subseteq S$ eine endliche Teilmenge. Definieren wir $T' = T \setminus E$, dann besitzt T die disjunkte Zerlegung $T = T' \cup (T \cap E)$, und außerdem gilt $T' \subseteq S \setminus E$. Auf Grund der Hilfsaussage existiert eine Teilmenge $F' \subseteq E \setminus S$ mit $|F'| = |T'|$ und der Eigenschaft, dass $(S \setminus T') \cup F'$ linear unabhängig ist. Sei nun $F = F' \cup (T \cap E)$. Wegen $T \subseteq S$ und $F' \cap S = \emptyset$ ist auch dies eine disjunkte Zerlegung, und folglich gilt $|F| = |F'| + |T \cap E| = |T'| + |T \cap E| = |T|$. Außerdem gilt

$$(S \setminus T) \cup F = ((S \setminus T') \setminus (T \cap E)) \cup F = (S \setminus T') \setminus (T \cap E) \cup (T \cap E) \cup F' = (S \setminus T') \cup F',$$

also ist diese Menge linear unabhängig. $\square$

Ein K -Vektorraum V wird **endlich erzeugt** genannt, wenn eine endliche Teilmenge $E \subseteq V$ mit $V = \langle E \rangle_K$ existiert.

(9.6) Satz Sei V ein endlich erzeugter K -Vektorraum.

- (i) In V existiert eine endliche Basis B .
- (ii) Für jede Basis B' von V gilt $|B'| = |B|$, insbesondere ist jede Basis endlich.
- (iii) Ist $S \subseteq V$ linear unabhängig und $E \supseteq S$ ein Erzeugendensystem von V , dann gibt es eine Basis B' von V mit $S \subseteq B' \subseteq E$.

Beweis: zu (i) Nach Voraussetzung existiert eine endliche Teilmenge $E_0 \subseteq V$ mit $V = \langle E_0 \rangle_K$. Weil E_0 endlich ist, existiert in E_0 eine maximale linear unabhängige Teilmenge B , die natürlich ebenfalls endlich ist. Nach Prop. (9.3) ist B also eine endliche Basis von V .

zu (ii) Sei B' eine weitere Basis von V . Zunächst zeigen wir, dass $|B'| \geq |B|$ gilt. Dazu wenden wir den Austauschsatz (9.5) auf $S = T = B$ und $E = B'$ an. Demzufolge existiert eine Teilmenge $F \subseteq B'$ mit $|F| = |B|$. Insbesondere gilt also die Ungleichung $|B'| \geq |F| = |B|$.

Um die Endlichkeit von B' zu beweisen, wählen wir in B' eine beliebige Teilmenge T' mit $|T'| = |B|$, was wegen $|B'| \geq |B|$ möglich ist. Der Austauschsatz liefert uns eine Teilmenge $F' \subseteq B$ mit der Eigenschaft, dass $|F'| = |T'| = |B|$ gilt und $(B' \setminus T') \cup F'$ linear unabhängig ist. Wegen $F' \subseteq B$ und $|F'| = |B|$ gilt $F' = B$. Folglich ist die Menge $(B' \setminus T') \cup B$ linear unabhängig. Weil aber B als Basis nach Satz (9.2) eine *maximale* linear unabhängige Teilmenge von V ist, muss $B' \setminus T' \subseteq B$ gelten. Es folgt $|B'| \leq |B' \setminus T'| + |T'| = |B' \setminus T'| + |B| \leq |B| + |B| = 2|B|$, insbesondere ist B' endlich.

Dasselbe Argument, dass oben die Ungleichung $|B| \leq |B'|$ gezeigt hat, kann nun auch auf B' an Stelle von B angewendet werden, und ergibt damit die Abschätzung $|B'| \leq |B|$. Insgesamt ist damit $|B'| = |B|$ gezeigt.

zu (iii) Sei $n = |B|$; wir zeigen zunächst, dass $|T| \leq n$ für jede linear unabhängige Menge mit $S \subseteq T \subseteq E$ gelten muss. Wenden wir den Austauschsatz auf T , eine beliebige *endliche* Teilmenge $T' \subseteq T$ und das Erzeugendensystem B an, so erhalten wir eine Teilmenge $F \subseteq B$ mit $|F| = |T'|$. Daraus folgt $n = |B| \geq |F| = |T'|$. Jede endliche Teilmenge von T hat also eine Mächtigkeit $\leq n$; dies ist nur möglich, wenn T endlich ist und $|T| \leq n$ gilt. Auf Grund der soeben bewiesenen Ungleichung gibt es in E eine *maximale* linear unabhängige (und endliche) Teilmenge B' mit $B' \supseteq S$. Aus Prop. (9.3) folgt, dass diese Teilmenge B' eine Basis von V ist. $\square$

Aus Teil (iii) von Satz (9.6) ergibt sich unmittelbar

(9.7) Folgerung Sei V ein endlich erzeugter K -Vektorraum.

- (i) (*Basisergänzungssatz*)
Jede linear unabhängige Teilmenge $S \subseteq V$ kann zu einer Basis von V ergänzt werden.
- (ii) (*Basisauswahlsatz*)
Aus jedem Erzeugendensystem E von V kann man eine Basis von V auswählen.

Beweis: Für Aussage (i) genügt es, Satz (9.6) (iii) auf S und $E = V$ anzuwenden. Für Aussage (ii) wendet man den Satz auf $S = \emptyset$ und E an. $\square$

(9.8) Definition Sei V ein K -Vektorraum. Dann ist die **Dimension** von V definiert durch

$$\dim V = \begin{cases} |B| & \text{falls } B \text{ eine endliche Basis von } V \text{ ist,} \\ \infty & \text{falls } V \text{ nicht endlich erzeugt ist.} \end{cases}$$

Man überprüfe anhand der bisherigen Resultate, dass die Dimension eines beliebigen K -Vektorraums damit wohldefiniert ist: Nach Satz (9.6) hat V entweder eine endliche Basis, oder V ist nicht endlich erzeugt. Im ersten Fall gilt außerdem für beliebig gewählte endliche Basen B, B' von V jeweils $|B| = |B'|$; es ist somit gleichgültig, welche Basis man für die Definition der Dimension heranzieht.

Wir bestimmen nun die Dimensionen der meisten uns bereits bekannten Vektorräume.

- (i) Ist K ein Körper und $n \in \mathbb{N}$, dann gilt $\dim K^n = n$. Denn wie wir bereits festgestellt haben, bilden die Einheitsvektoren $e_1, \dots, e_n$ in K^n eine n -elementige Basis.
- (ii) Seien K ein Körper und $m, n \in \mathbb{N}$. Dann gilt $\dim \mathcal{M}_{m \times n, K} = mn$, weil die Basismatrizen $B_{kl}^{(m \times n)}$ mit $1 \leq k \leq m$ und $1 \leq l \leq n$ eine mn -elementige Basis von $\mathcal{M}_{m \times n, K}$ bilden.
- (iii) Wir wissen bereits, dass $\mathbb{C}$ auf natürliche Weise als $\mathbb{R}$ -Vektorraum angesehen werden kann. Eine Basis dieses Vektorraums ist durch $\{1, i\}$ gegeben. Denn einerseits kann jedes $z \in \mathbb{C}$ auf Grund der Zerlegung in Real- und Imaginärteil in der Form $z = a \cdot 1 + b \cdot i$ mit $a, b \in \mathbb{R}$ dargestellt werden. Dies zeigt, dass $\{1, i\}$ ein Erzeugendensystem ist. Andererseits ist diese Darstellung auch eindeutig, denn aus $z = a \cdot 1 + b \cdot i$ mit $a, b \in \mathbb{R}$ folgt $a = \operatorname{Re}(z)$ und $b = \operatorname{Im}(z)$. Deshalb ist $\{1, i\}$ auch linear unabhängig. Für $\mathbb{C}$ als $\mathbb{R}$ -Vektorraum gilt also $\dim \mathbb{C} = 2$; um zu verdeutlichen, dass $\mathbb{C}$ als $\mathbb{R}$ -Vektorraum betrachtet wird, schreibt man $\dim_{\mathbb{R}} \mathbb{C} = 2$.
- (iv) Fassen wir $\mathbb{C}$ dagegen als $\mathbb{C}$ -Vektorraum auf, dann ist $\{1\}$ eine Basis, und es gilt $\dim \mathbb{C} = \dim_{\mathbb{C}} \mathbb{C} = 1$. Allgemein ist es nicht schwer zu zeigen, dass $\dim_{\mathbb{C}} \mathbb{C}^n = n$ und $\dim_{\mathbb{R}} \mathbb{C}^n = 2n$ für alle $n \in \mathbb{N}$ gilt; eventuell erledigen wir das in den Übungen.
- (v) Ist K ein Körper und V ein K -Vektorraum mit $V = \{0_V\}$, dann gilt $\dim V = 0$, denn in diesem Fall ist die leere Menge $\emptyset$ eine nullelementige Basis von V . Tatsächlich ist $\emptyset$ linear unabhängig, und die einzige Linearkombination von $\emptyset$ ist der Nullvektor; es gilt also $\langle \emptyset \rangle_K = \{0_V\}$. Man kann sich leicht überlegen, dass umgekehrt aus $\dim V = 0$ jeweils $V = \{0_V\}$ folgt.
- (vi) Für jeden Körper K gilt $\dim K[x] = \infty$. Denn wie wir in § 6 festgestellt haben, besitzt $K[x]$ als K -Vektorraum kein endliches Erzeugendensystem.

Als weitere Konsequenz aus Satz (9.6) notieren wir noch

(9.9) Folgerung Sei V ein endlich erzeugter K -Vektorraum und $n = \dim V$.

- (i) Für jede linear unabhängige Teilmenge $S \subseteq V$ gilt $|S| \leq n$ mit Gleichheit genau dann, wenn S eine Basis von V ist.
- (ii) Für jedes Erzeugendensystem E von V gilt $|E| \geq n$ mit Gleichheit genau dann, wenn E eine Basis von V ist.

Beweis: zu (i) Sei $S \subseteq V$ linear unabhängig. Nach dem Basisergänzungssatz gibt es eine Basis B von V mit $B \supseteq S$. Daraus folgt $|S| \leq |B| = \dim V = n$. Beweisen wir nun die Äquivalenz. Gilt $|S| = n$, dann folgt aus $S \subseteq B$ und $|S| = n = |B|$ die Gleichheit $S = B$. In diesem Fall ist S also selbst eine Basis. Setzen wir umgekehrt voraus, dass S eine Basis von V ist, dann muss $|S| = \dim V = n$ gelten, denn die Dimension von V kann mit jeder beliebigen Basis bestimmt werden.

zu (ii) Sei $E \subseteq V$ ein Erzeugendensystem. Nach dem Basisauwahlsatz gibt es eine Basis $B \subseteq E$ von V . Daraus folgt $n = \dim V = |B| \leq |E|$. Gilt $|E| = n$ dann folgt aus $|E| = n = |B|$ und $B \subseteq E$ die Gleichheit $E = B$. Also ist E in diesem Fall selbst eine Basis. Setzen wir andererseits voraus, dass E eine Basis von V ist, dann muss $|E| = \dim V = n$ gelten, mit demselben Argument wie in Teil (i). $\square$

Die soeben bewiesene Aussage kann folgendermaßen praktisch genutzt werden: Wenn man von einem endlich erzeugten K -Vektorraum V die (endliche) Dimension n bereits kennt und $T \subseteq V$ eine n -elementige Teilmenge ist, dann folgt aus der linearen Unabhängigkeit bereits die Basiseigenschaft von T . Ebenso folgt aus $V = \langle T \rangle_K$ bereits die Basiseigenschaft.

Beispielsweise ist wegen $\dim \mathbb{R}^3 = 3$ jede dreielementige linear unabhängige Teilmenge des $\mathbb{R}^3$ bereits eine Basis von $\mathbb{R}^3$, und ebenso ist jedes dreielementige Erzeugendensystem eine Basis. Andererseits zeigt die Folgerung auch, dass es in $\mathbb{R}^3$ kein zweielementiges Erzeugendensystem und keine vierelementige linear unabhängige Teilmenge geben kann.

(9.10) Folgerung Sei V ein endlich erzeugter K -Vektorraum, $n = \dim V$ und U ein Untervektorraum von V . Dann gilt $\dim U \leq n$ mit Gleichheit genau dann, wenn $U = V$ gilt.

Beweis: Sei B eine Basis von U . Dann ist B insbesondere eine linear unabhängige Teilmenge von V , und aus Folgerung (9.9) (i) erhalten wir $\dim U = |B| \leq n$. Setzen wir $U = V$ voraus, dann folgt offenbar $\dim U = \dim V = n$. Sei nun umgekehrt $\dim U = n = \dim V$ vorausgesetzt. Dann ist B wegen $|B| = \dim U$ eine n -elementige linear unabhängige Teilmenge von V . Aus Folgerung (9.9) (i) ergibt sich, dass B eine Basis von V ist. Somit gilt $U = \langle B \rangle_K = V$. $\square$

Zum Abschluss des Kapitels soll eine praktische Umsetzung von Basisauswahlsatz und Basisergänzungssatz diskutiert werden. Konkret beantworten wir die folgenden beiden Fragen. Sei V ein endlich-dimensionaler K -Vektorraum.

- Wenn $E = \{v_1, \dots, v_r\}$ ein Erzeugendensystem von V ist, wie findet man eine in E enthaltene Basis?
- Wenn $S = \{v_1, \dots, v_r\}$ eine r -elementige linear unabhängige Teilmenge von V ist, wie lässt sich die Menge S zu einer Basis von V ergänzen?

Der folgende Satz liefert eine Antwort auf die erste Frage im Spezialfall $V = K^m$. Wir werden später sehen, dass jeder endlich-dimensionale K -Vektorraum $V \neq \{0_V\}$ isomorph zu K^m für ein $m \in \mathbb{N}$ ist. Das Problem der Basisauswahl lässt sich dann leicht auf solche Vektorräume übertragen.

(9.11) Satz Seien $m, n \in \mathbb{N}$, und sei $S = \{v_1, \dots, v_n\}$ eine n -elementige Teilmenge von K^m . Sei $A \in \mathcal{M}_{m \times n, K}$ die Matrix, deren Spalten genau die Vektoren $v_1, \dots, v_n$ sind, und sei A' die Matrix in normierter ZSF, die man durch Anwendung des Gauß-Verfahrens auf A erhält. Seien $r, j_1, \dots, j_r$ die Kennzahlen der ZSF. Dann ist $\{v_{j_1}, \dots, v_{j_r}\}$ eine Basis von $\langle S \rangle_K$.

Beweis: Die gemeinsame Lösungsmenge $\mathcal{L} \subseteq K^n$ der homogenen linearen Gleichungssysteme mit den Koeffizientenmatrizen A und A' sind genau die Tupel $(\lambda_1, \dots, \lambda_n)$ mit der Eigenschaft $\lambda_1 v_1 + \dots + \lambda_n v_n = 0_{K^m}$. Sei $S = \{1, \dots, n\} \setminus \{j_1, \dots, j_r\}$. Wir betrachten nun das LGS zur Matrix A' . Bilden wir für ein beliebiges $\ell \in S$ den Lösungsvektor b_ℓ wie es im Kapitel zur Lösung von LGS beschrieben, dann enthält dieser an der ℓ -ten Position den Wert 1, und die übrigen Einträge ungleich null müssen sich auf die Positionen $j_1, \dots, j_r$ verteilen. Dies zeigt, dass der Vektor v_ℓ eine Linearkombination von $T = \{v_{j_1}, \dots, v_{j_r}\}$ ist. Insgesamt zeigt dies, dass die Vektoren $v_1, \dots, v_n$ alle in $\langle T \rangle_K$ enthalten und T folglich ein Erzeugendensystem von $\langle S \rangle_K$ ist.

Wäre T keine Basis von $\langle S \rangle_K$, also linear abhängig, dann müsste es nach Prop. (8.9) (i) möglich sein, ein Element v_{j_k} als Linearkombination der Vektoren v_{j_s} mit $s \neq k$ darzustellen. Es gäbe dann in $\mathcal{L}$ ein Element der Form $(\lambda_1, \dots, \lambda_n)$ mit $\lambda_{j_k} = 1$ und $\lambda_\ell = 0$ für alle $\ell \in S$. Setzt man diese Werte aber in die der k -ten Zeile von A' entsprechenden Gleichung ein, in der (nach Definition der normierten ZSF) alle Koeffizienten der Variablen x_{j_s} mit $s \neq k$ gleich null sind, so erhält man die falsche Gleichung $1 = 0$. Der Widerspruch zeigt, dass es in $\mathcal{L}$ kein derartiges Element und somit auch keine Darstellung von v_{j_k} als Linearkombination der Vektoren v_{j_s} mit $s \neq k$ existiert. $\square$

Wir demonstrieren die Anwendung des Satzes an einem konkreten Beispiel. Unser Ziel ist es, aus der Teilmenge $S = \{v_1, v_2, v_3\}$ des $\mathbb{R}^3$ gegeben durch

$$v_1 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}, \quad v_2 = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \quad \text{und} \quad v_3 = \begin{pmatrix} 2 \\ 3 \\ 4 \end{pmatrix}$$

eine Basis des Vektorraums $V = \langle S \rangle_K$ auszuwählen. Dazu tragen wir die Vektoren als Spalten in eine Matrix ein und formen auf normierte ZSF um.

$$\begin{pmatrix} 1 & 1 & 2 \\ 2 & 1 & 3 \\ 3 & 1 & 4 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 1 & 2 \\ 0 & -1 & -1 \\ 0 & -2 & -2 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 1 & 2 \\ 0 & 1 & 1 \\ 0 & -2 & -2 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 1 & 2 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 0 \end{pmatrix}$$

Die normierte ZSF am Ende hat die Kennzahlen $r = 2$, $j_1 = 1$ und $j_2 = 2$. Nach Satz (9.11) ist somit $\{v_1, v_2\}$ eine Basis von $\langle S \rangle_K$. Anhand der Lösungsmenge des zur Matrix gehörenden homogenen LGS lässt sich auch leicht erkennen, dass der Vektor v_3 als Linearkombination von $\{v_1, v_2\}$ dargestellt werden kann und somit für eine Basis von $\langle S \rangle_K$ nicht benötigt wird. Die Matrix in normierter ZSF entspricht dem LGS bestehend aus den Gleichungen $x_1 + x_3 = 0$, $x_2 + x_3 = 0$. Die Lösungsmenge ist somit gegeben durch

$$\mathcal{L} = \left\{ \begin{pmatrix} x_1 \\ x_2 \\ x_3 \end{pmatrix} \in \mathbb{R}^3 \mid x_1 = -x_3, x_2 = -x_3 \right\} = \left\{ \begin{pmatrix} -x_3 \\ -x_3 \\ x_3 \end{pmatrix} \mid x_3 \in \mathbb{R} \right\} = \left\{ x_3 \begin{pmatrix} -1 \\ -1 \\ 1 \end{pmatrix} \mid x_3 \in \mathbb{R} \right\}.$$

Der Lösungsvektor $(-1, -1, 1) \in \mathcal{L}$ liefert die Gleichung $(-1)v_1 + (-1)v_2 + v_3 = 0_{\mathbb{R}^3}$, was zu $v_3 = v_1 + v_2$ äquivalent ist.

Kommen wir nun zu unserer zweiten Frage. Gegeben sei eine linear unabhängige Teilmenge $S = \{v_1, \dots, v_n\}$ im K^m , die zu einer Basis von K^m ergänzt werden soll. Wir wissen bereits, dass die Menge $\{e_1, \dots, e_m\}$ der Einheitsvektoren eine Basis und damit erst recht ein Erzeugendensystem des K^m bildet. Also ist auch die Menge $T = \{v_1, \dots, v_n, e_1, \dots, e_m\}$ ein Erzeugendensystem des K^m . Mit dem in Satz (9.11) formulierten Kriterium kann aus T eine Basis ausgewählt werden. Dabei ist nur zu beachten, dass die Vektoren $v_1, \dots, v_n, e_1, \dots, e_m$ tatsächlich in dieser Reihenfolge als Spalten in die Matrix A eingetragen werden.

Für $1 \leq \ell \leq n+m$ wird der ℓ -te Vektor der Menge T vom Algorithmus genau dann aus dem Erzeugendensystem entfernt, wenn ℓ in der Menge S liegt, also nicht unter den Kennzahlen $j_1, \dots, j_r$ der normierten ZSF vorkommt. Der zugehörige Lösungsvektor b_ℓ besitzt Einträge ungleich null nur an den Stellen j_k mit $k < \ell$ (denn nach Definition der normierten ZSF kann der Eintrag $a'_{k\ell}$ der ℓ -ten Spalte nur dann ungleich null sein, wenn $j_k < \ell$ ist). Dies bedeutet, dass der Vektor v_ℓ eine Linearkombination der Vektoren v_{j_k} mit $j_k < \ell$ ist. Weil die Menge $S = \{v_1, \dots, v_n\}$ linear unabhängig ist, ist kein v_k mit $k \in \{1, \dots, n\}$ als Linearkombination von $v_1, \dots, v_{k-1}$ darstellbar. Dies bedeutet, dass keiner der Vektoren v_k aus dem Erzeugendensystem entfernt wird. Somit ist gewährleistet, dass wir tatsächlich eine Basis von K^m erhalten, die S als Teilmenge enthält.

Auch die Basisergänzung demonstrieren wir an einem konkreten Beispiel. Wie man mit Proposition (8.9) (i) leicht überprüft, ist die Menge $S = \{v_1, v_2\}$ bestehend aus den Vektoren

$$v_1 = \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix} \quad \text{und} \quad v_2 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix}$$

linear unabhängig. Unser Ziel besteht darin, S zu einer Basis von $\mathbb{R}^3$ zu ergänzen. Dazu schreiben wir die Vektoren v_1, v_2, e_1, e_2, e_3 als Spalten in eine Matrix und formen diese auf normierte ZSF um.

$$\begin{pmatrix} 2 & 1 & 1 & 0 & 0 \\ 2 & 2 & 0 & 1 & 0 \\ 3 & 3 & 0 & 0 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{2} & 0 & 0 \\ 2 & 2 & 0 & 1 & 0 \\ 3 & 3 & 0 & 0 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{2} & 0 & 0 \\ 0 & 1 & -1 & 1 & 0 \\ 0 & \frac{3}{2} & -\frac{3}{2} & 0 & 1 \end{pmatrix} \mapsto \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{2} & 0 & 0 \\ 0 & 1 & -1 & 1 & 0 \\ 0 & 0 & 0 & -\frac{3}{2} & 1 \end{pmatrix}$$

$$\mapsto \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{2} & 0 & 0 \\ 0 & 1 & -1 & 1 & 0 \\ 0 & 0 & 0 & 1 & -\frac{2}{3} \end{pmatrix} \mapsto \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{2} & 0 & 0 \\ 0 & 1 & -1 & 0 & \frac{2}{3} \\ 0 & 0 & 0 & 1 & -\frac{2}{3} \end{pmatrix} \mapsto \begin{pmatrix} 1 & 0 & 1 & 0 & -\frac{1}{3} \\ 0 & 1 & -1 & 0 & \frac{2}{3} \\ 0 & 0 & 0 & 1 & -\frac{2}{3} \end{pmatrix}$$

Die normierte ZSF hat die Kennzahlen $r = 3$, $j_1 = 1$, $j_2 = 2$, $j_3 = 4$. Mit Satz (9.11) folgt daraus, dass $B = \{v_1, v_2, e_2\}$ eine Basis von $\mathbb{R}^3$ ist, die zudem S als Teilmenge enthält.

Ähnlich wie im vorherigen Beispiel findet man durch Bestimmung der Lösungsmenge $\mathcal{L} \subseteq \mathbb{R}^5$ des homogenen LGS zur umgeformten Matrix konkrete Darstellungen von e_1 und e_3 als Linearkombinationen der Basis B ; es gilt

$$1 \cdot \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix} + (-1) \cdot \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$$

und

$$\left(-\frac{1}{3}\right) \cdot \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix} + \frac{2}{3} \cdot \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} + \left(-\frac{2}{3}\right) \cdot \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}.$$

Literaturverzeichnis

[Bo] S. Bosch, *Lineare Algebra*. Springer-Lehrbuch, Berlin 2006.

[dJ] T. de Jong, *Lineare Algebra*. Pearson-Studium, München 2013.

[Fi] G. Fischer, *Lernbuch Lineare Algebra und Geometrie*. Vieweg-Teubner, Wiesbaden 2011.

[Jn] K. Jaenich, *Lineare Algebra*. Springer-Verlag, Berlin 2001.